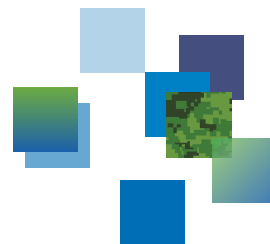




DRDC | RDDC



Bayesian inference for source term estimation: Application to the International Monitoring System radionuclide network

Eugene Yee

DRDC – Suffield Research Centre

Ian Hoffman

Radiation Protection Branch, Health Canada

Kurt Ungar

Radiation Protection Branch, Health Canada

Alain Malo

Canadian Meteorological Centre, Environment Canada

Nils Ek

Canadian Meteorological Centre, Environment Canada

Pierre Bourgouin

Canadian Meteorological Centre, Environment Canada

Defence Research and Development Canada

Scientific Report

DRDC-RDDC-2014-R71

October 2014

Bayesian inference for source term estimation: Application to the International Monitoring System radionuclide network

Eugene Yee

DRDC – Suffield Research Centre

Ian Hoffman

Radiation Protection Branch, Health Canada

Kurt Ungar

Radiation Protection Branch, Health Canada

Alain Malo

Canadian Meteorological Centre, Environment Canada

Nils Ek

Canadian Meteorological Centre, Environment Canada

Pierre Bourgouin

Canadian Meteorological Centre, Environment Canada

Defence Research and Development Canada

Scientific Report

DRDC-RDDC-2014-R71

October 2014

- © Her Majesty the Queen in Right of Canada, as represented by the Minister of National Defence, 2014
- © Sa Majesté la Reine en droit du Canada, telle que représentée par le ministre de la Défense nationale, 2014

Abstract

In recent years, there has been an enormous quantity of data obtained from the International Monitoring System radionuclide network for the verification of the Comprehensive Nuclear-Test-Ban Treaty. The complexity of the instruments deployed here, of the radionuclide sources, and of the myriad of scientific questions related to treaty verification lead invariably to complex inference problems (associated with source term estimation) that require the application of sophisticated statistical tools. In this report, we demonstrate that a rigorous and general framework for addressing these problems is through Bayesian probability theory, allowing the rational inference of the posterior probability distribution of the source parameters of interest given any prior information and available activity concentration measurements. The methodology is demonstrated by application to two different problems: namely, the emission rate profile reconstruction of a radioxenon release from the Fukushima Daiichi nuclear power plant and source reconstruction (location and emission rate) of a radioxenon release from the Chalk River Laboratories (CRL) medical isotope production facility. The sampling of the resulting posterior distribution of the source parameters is undertaken using two different Markov chain Monte Carlo techniques: namely, nested sampling and multiple-try differential evolution adaptive Metropolis sampling with a past archive.

For the Fukushima nuclear power plant release, it is demonstrated that the *limited* temporal extent of the activity concentration time series obtained from the seven sampling sites cannot be used to constrain the emission rate profile at a later time. In particular, the Bayesian credible intervals for the reconstruction of the emission rate profile provide a quantitative indication of the uncertainty in this quantity, allowing an objective assessment of the fact that the emission rates recovered at the later times are not constrained by the information in the available activity concentration data. For the CRL release, it is shown that the principal difficulty in the reconstruction lay in the correct specification of both the scale and structure of the model errors used in the Bayesian inferential methodology. Consequently, for this case, two different measurement models for incorporation of the model errors in the predicted concentrations are considered. The performance of both of these measurement models with respect to their accuracy and precision in the recovery of the source parameters is compared and contrasted. In particular, it is shown that the incorporation of multipliers in the measurement model to compensate for the unknown model errors lead to significantly improved estimates for the source parameter (both in accuracy and in precision) compared to the simpler measurement model which does not use multipliers.

Significance for defence and security

The International Monitoring System (IMS) radionuclide network has been established for the verification of the Comprehensive Nuclear-Test-Ban Treaty (CTBT). Compliance with the CTBT is a global security issue, and the work reported herein is significant in that it provides novel and advanced methods for model-based statistical inference relevant to the detection and rapid estimation of an unknown source term associated with a covert nuclear test. These model-based methods for statistical inference provide a potentially new approach for a sensor-driven modeling paradigm involving the fusion of sensor measurements of radionuclide activity concentration with the predictive outputs (model activity concentration) obtained from advanced models for atmospheric dispersion. This may lead potentially to significant improvements in situational awareness in the ‘battlespace’ with respect to the important problem of the determination of the characteristics of an unknown source resulting from a clandestine nuclear test following the detection of the event by the IMS radionuclide network.

Résumé

Au cours des dernières années, des quantités considérables de données ont été recueillies grâce au réseau du système de surveillance international (SSI) des radionucléides pour la vérification du Traité d'interdiction complète des essais nucléaires (TICEN). La complexité des instruments déployés, des sources de radionucléides et des innombrables questions scientifiques relatives à la vérification du traité donnent inmanquablement lieu à des problèmes de déduction complexes (associés à l'estimation du terme source), ce qui nécessite l'utilisation d'outils statistiques perfectionnés. Dans le présent rapport, nous allons démontrer qu'un cadre rigoureux et général pour s'attaquer à ces problèmes est offert par la théorie des probabilités de Bayes et que celui-ci permet de procéder à une déduction rationnelle de la distribution de probabilités ultérieure des paramètres des sources, peu importe les informations et mesures de l'activité volumique disponibles obtenues antérieurement. La méthode est démontrée par son application à deux problèmes différents : la reconstruction du profil du taux d'émission relatif à un rejet de xénon radioactif par la centrale nucléaire de Fukushima Daiichi et la reconstruction des sources (emplacement et taux d'émission) d'un rejet de xénon radioactif à l'installation de production d'isotopes médicaux des Laboratoires de Chalk River (LCR). L'échantillonnage de la distribution ultérieure des paramètres des sources qui en résulte est réalisé à l'aide de deux techniques différentes de simulation Monte Carlo par chaînes de Markov, soit l'échantillonnage à emboîtements et l'échantillonnage Metropolis adaptatif à évolution différentielle et à essais multiples, utilisant des archives antérieures.

Dans le cas du rejet de la centrale nucléaire de Fukushima, il est démontré que la portée temporelle *limitée* de la série chronologique de l'activité volumique obtenue à partir des sept sites d'échantillonnage ne peut pas être utilisée pour restreindre le profil du taux d'émission ultérieurement. Plus particulièrement, les intervalles de crédibilité de Bayes pour la reconstruction du profil du taux d'émission fournissent une indication quantitative de l'incertitude à cette quantité, ce qui permet une évaluation objective du fait que les taux d'émission récupérés à des dates ultérieures ne sont pas restreints par l'information provenant des données disponibles sur l'activité volumique. Dans le cas des rejets des LCR, il est indiqué que la principale difficulté de la reconstruction est la spécification correcte de l'échelle et de la structure des erreurs du modèle utilisé dans la méthode par déduction de Bayes. Par conséquent, dans ce cas particulier, deux modèles de mesure différents sont envisagés pour l'intégration des erreurs du modèle dans aux concentrations prévues. Les rendements de ces deux modèles de mesure, en ce qui a trait à leur exactitude et à leur précision dans la récupération des paramètres des sources sont comparés et leurs éléments importants sont mis en évidence. Plus particulièrement, on montre que l'intégration des multiplicateurs au modèle de mesure dans le but de compenser les erreurs de modèle inconnues a donné lieu à des estimations considérablement améliorées des paramètres

des sources (à la fois sur le plan de l'exactitude et de la précision), comparativement au modèle de mesure plus simple qui n'utilise pas de multiplicateurs.

Importance pour la défense et la sécurité

Le réseau du système de surveillance international (SSI) des radionucléides (IMS, de l'anglais *International Monitoring System*) a été mis sur pied pour vérifier le Traité d'interdiction complète des essais nucléaires (TICEN). La conformité au TICEN est une question de sécurité mondiale, et les travaux consignés ici sont importants en ce sens qu'ils offrent des méthodes nouvelles et avancées pour la déduction statistique basée sur un modèle concernant la détection et l'estimation rapide d'un terme source inconnu associé à des essais nucléaires clandestins. Ces méthodes de déduction statistique basées sur un modèle fournissent une nouvelle démarche possible au paradigme de modélisation déterminée par des capteurs, fusionnant les mesures de capteurs d'activité volumique de radionucléides et les prévisions (activité volumique du modèle) obtenues à partir de modèles avancés de dispersion atmosphérique. Cela pourrait améliorer considérablement la connaissance de la situation dans l'« espace de bataille » en ce qui a trait au problème important de la détermination des caractéristiques d'une source inconnue résultant d'un essai nucléaire clandestin à la suite de la détection de l'activité par le réseau du SSI des radionucléides.

Table of contents

Abstract	i
Significance for defence and security	ii
Résumé	iii
Importance pour la défense et la sécurité	iv
Table of contents	v
List of figures	vi
List of tables	viii
1 Introduction	1
2 Bayesian inference in a nutshell	4
3 Assignment of the direct probabilities	7
3.1 Assignment of likelihood – input	7
3.2 Assignment of prior – input	10
3.3 Posterior distribution – output	10
4 Bayesian computation	12
4.1 Rapid computation of model concentration	12
4.2 Markov chain Monte Carlo	14
4.2.1 Differential evolution adaptive Metropolis sampling	15
4.2.2 Nested sampling	17
5 Applications	19
5.1 Case 1: Emission rate profile reconstruction	20
5.2 Case 2: Source reconstruction	26
6 Discussion and conclusions	35
References	39

List of figures

Figure 1:	Locations of the seven sampling stations (yellow markers) from the International Monitoring System radionuclide network used for the emission rate profile reconstruction. The location of the Xe-133 tracer source was assumed to be known <i>a priori</i> to be at the Fukushima Daiichi nuclear power plant (red marker).	20
Figure 2:	The best estimate for the Xe-133 emission rate time profile given by the posterior mean (solid red curve) compared to the true emission rate time profile (solid black curve), along with the 90% credible (HPD) interval for the emission rate profile estimate (gray-shaded region).	23
Figure 3:	Marginal posterior probability distribution for $\log_{10}(Q)$ obtained for the 12-h time interval (a) (48,60) h and (b) (168,180) h after the arbitrary time origin t_0 (11 March 2011 00:00 UTC). The solid vertical line indicates the true value for the emission rate and the dashed vertical line corresponds to the best estimate of the emission rate obtained as the posterior mean of the associated marginal posterior probability distribution.	24
Figure 4:	Locations of the three sampling stations (yellow markers) from the International Monitoring System radionuclide network used for Case 2. The location of the Xe-133 tracer source (red marker) was at Chalk River Laboratories.	27
Figure 5:	Univariate (diagonal) and bivariate (off-diagonal) marginal posterior distributions of the source parameters: longitudinal position x_s , latitudinal position y_s , and emission rate Q_s . A solid square or a solid vertical line indicates the true parameter values, whereas a solid circle or a dashed vertical line indicates a best estimate (posterior mean) for these parameter values.	28
Figure 6:	Two-dimensional marginal posterior distribution of the source location geo-referenced on a Google Earth image.	30

Figure 7:	Univariate (diagonal) and bivariate (off-diagonal) marginal posterior distributions of the source parameters: longitudinal position x_s , latitudinal position y_s , and emission rate Q_s . The true parameter values are shown by a solid square or a solid vertical line and the best estimates of the parameter values are represented as a solid circle or a dashed vertical line. The reconstruction was obtained using the alternative measurement model.	31
Figure 8:	Two-dimensional marginal posterior distribution of the source location geo-referenced on a Google Earth image. The marginal distribution is for the alternative measurement model.	33
Figure 9:	Univariate marginal posterior distribution of the multipliers m_J ($J = 1, 2, \dots, 8$). The “actual” values of the multipliers are indicated with the solid vertical line which should be compared with the best estimates (posterior means) of the multipliers. . . .	34

List of tables

Table 1:	The latitudinal and longitudinal coordinates of the seven sampling stations used for the emission rate profile reconstruction.	22
Table 2:	The posterior mean, posterior standard deviation, and lower and upper bounds of the 95% HPD interval of the parameters x_s ($^{\circ}\text{E}$), y_s ($^{\circ}\text{N}$), and Q_s ($\mu\text{Bq h}^{-1}$) calculated from samples of $\boldsymbol{\theta}$ drawn from the posterior distribution $p(\boldsymbol{\theta} \mathbf{d}, \mathbf{s}, \alpha, \beta, I)$	29
Table 3:	The posterior mean, posterior standard deviation, and lower and upper bounds of the 95% HPD interval of the parameters x_s ($^{\circ}\text{E}$), y_s ($^{\circ}\text{N}$), and Q_s ($\mu\text{Bq h}^{-1}$) calculated from samples of $\boldsymbol{\theta}$ drawn from the posterior distribution $p(\boldsymbol{\theta} \mathbf{d}, \mathbf{s}, \alpha, \beta, I)$. The results are obtained using the alternative measurement model.	32

1 Introduction

The monitoring of noble gas radionuclides and, in particular, xenon radionuclides (radioxenon) is an important component to ensure compliance with the Comprehensive Nuclear-Test-Ban Treaty (CTBT) through the prompt detection of clandestine nuclear testing (whether underground, underwater, or in the atmosphere). This is because radioxenon is by far the most abundant of the noble gas radionuclides released during a nuclear detonation, making the monitoring and measurement of radioxenon critically important for the verification or dismissal of a putative covert nuclear test (especially for the case of an underground or underwater test where the presence of particulate radionuclides is minimized).

One of the main challenges with radionuclide monitoring for CTBT verification concerns the problem of source term estimation or reconstruction (viz., determination of unknown source characteristics such as location, emission rate, and time of release following event detection). This problem is particularly acute for noble gas monitoring for CTBT verification because xenon radionuclides can enter the atmosphere from various other sources such as the normal operation of nuclear reactors and the production and use of medical isotopes [1] which can create difficulties in the interpretation of the measurement results for source attribution (viz., monitoring for compliance with the CTBT).

Noble gas monitoring of radioxenon for Comprehensive Nuclear-Test-Ban Treaty purposes was primarily envisioned for the detection of underground nuclear tests, but it can also be applied to detect nuclear tests that occur in the atmosphere. The most significant problem in noble gas monitoring is the discrimination of sources. Owing to the production of medical isotopes, an almost global radioxenon background (composed principally of ^{133}Xe , but other isomers such as $^{133\text{m}}\text{Xe}$ and $^{131\text{m}}\text{Xe}$ are also important) is created that must be decoupled from any nuclear test signal of interest. Medical isotope production has been demonstrated to perturb measurements that occur thousands of kilometers distant from the production site [2]. Therefore, to perform an accurate analysis of the genesis time using nuclide ratios, some compensation is required to account for these medical isotope emissions. The genesis time, corrected for background, is extremely important for the source localization problem [1].

Currently, typical source term estimation approaches use backward Lagrangian particle transport and dispersion modeling in combination with re-analyzed global wind fields to determine source-receptor sensitivity (SRS) fields (or, retro-plume concentration fields) for each measurement [3]. The SRS fields provide information on the sensitivity relationship between the emission rate at a given source space-time point and the activity concentration measured at the receptor. The SRS fields can be used to establish a field of regard (FOR) which is simply the geographical region that is covered by the SRS fields for a particular activity concentration measurement over a

fixed (prescribed) time interval preceding this measurement.

In the case of multiple (related) activity concentration measurements, a separate FOR can be generated for each measurement and then combined to create a possible source region (PSR). There are two approaches that have been used to combine multiple measurements together into a PSR to provide a common geo-temporal area where the source may be located. Wotawa et al. [3] used a spatial-temporal correlation distribution between the modeled sensitivity and the concentration measurements to provide a measure (or, score) of how well a specified release at a given location and time would be able to explain the behavior of the available measurements. Another approach was proposed by Ringbom et al. [4] in which the separate FORs were superimposed to create a PSR containing only the geographical area that is common to all the individual FOR regions (viz., the common intersection between all the FORs). This was the approach used in the analysis of the most recently announced Democratic People's Republic of Korea nuclear test [4]. This geographically and temporally consistent region is more restrictive in spatial-temporal extent than that obtained using the correlation distribution methodology as advocated by Wotawa et al. [3].

Further reductions in the geographical area associated with the PSR can be achieved through the inclusion of additional information in the analysis. For example, the area enclosing the putative source can be reduced by the utilization of the timing information derived from nuclide activity ratios (if this is available), or by incorporating the information regarding non-detections from the monitoring network into the analysis process. The ideal information for source localization would be the measurement of an unambiguous seismic signal. However, even here, problems can arise when the received seismic signals are near the discrimination threshold, as then it becomes increasingly difficult to separate natural seismic events from potential explosive (nuclear test) events. Furthermore, it should be noted that while a seismic signal may provide evidence that an event has occurred, only the detection of radionuclides (or other by-products released by a nuclear detonation) can provide the definitive verification that the putative event was the result of a nuclear explosion.

The methods that have been used for source term reconstruction involving the manipulation of the source-receptor sensitivity field(s) to provide either a FOR or a PSR are *deterministic* in nature in the sense that they select a single source distribution from the entire ensemble of acceptable distributions that are consistent with the finite number of noisy activity concentration data provided by the sensor network. As such, the source term estimation using these methodologies does not provide information on the reliability (or, equivalently, uncertainty) of the solution. This report focuses on the application of a novel approach for source term estimation which is based on a *probabilistic* approach using a Bayesian inferential scheme that allows the uncertainty in the inference of the source characteristics to be rigorously quantified. More specifically, this proposed alternative to the existing approaches described above in-

volves the characterization of an ensemble of source distributions using the Bayesian statistical paradigm (providing a fully probabilistic solution allowing the uncertainty in the source reconstruction to be rigorously assessed). In this report, we provide an argument of *why* it is necessary to apply Bayesian probability theory and a detailed description on *how* we should use it for source reconstruction.

The Bayesian inferential methodology for source reconstruction provides fully probabilistic information on all the parameters used to describe the unknown source distribution. The probabilistic approach using a Bayesian inferential scheme for source reconstruction has been developed (in other source reconstruction contexts). The methodology was utilized by Yee [5] and demonstrated initially using Project Prairie Grass data for short-range dispersion over open and level terrain. The Bayesian inferential scheme for source term reconstruction was further developed, refined and generalized in subsequent work which included (1) application of the methodology to complex environments (including dispersion in urban environments) by Yee [6], Keats et al. [7] and Chow et al. [8]; (2) generalization of the methodology to deal with a non-conservative scalar by Keats et al. [9]; (3) Bayesian experimental design for receptor placement in order to maximize the expected information in the measured concentration data for improving estimates of the source location and emission rate [10]; (4) application of the methodology to source reconstruction for long-range dispersion on continental scales by Yee et al. [11]; and, (5) reconstruction of multiple sources when the number of sources is unknown *a priori* was addressed by Yee [12, 13, 14] and by Yee and Flesch [15] as a generalized parameter estimation problem and by Yee [16] as a model selection problem.

Most of the applications of the Bayesian inferential methodology for source reconstruction have used high-quality concentration data from well-designed atmospheric transport and dispersion experiments to validate the schema. The objective of this report is to use activity concentration data (synthetic and real) obtained from an operational network of sensors (more specifically, from a very small subset of the global network of radionuclide sensors that form part of the International Monitoring System deployed under the auspices of the Comprehensive Nuclear-Test-Ban Treaty [17]) to provide a real-world test of source reconstruction based on the Bayesian statistical paradigm applied to long-range atmospheric transport on a continental or hemispheric scale. A preliminary description of the application of Bayesian inference to source reconstruction using measurements obtained from some radionuclide sensors of the International Monitoring System has been reported recently by Yee et al. [18]. In this report, we will give a much more comprehensive and detailed description of the application of Bayesian inference to real-world applications for source term reconstruction involving activity concentration data obtained from the International Monitoring System. The purpose is to provide the reader with a review summarizing the general ideas of Bayesian inference as applied to source reconstruction along with an overview of numerical techniques that are useful for Bayesian analysis and to

illustrate the application of these ideas and numerical techniques using two relevant case studies.

It will be demonstrated that one of the key problems that needs to be properly addressed for the successful application of Bayesian inference for source reconstruction when using operational concentration data obtained from the International Monitoring System is the ‘correct’ specification of the expected model errors arising from the putative accuracy of a long-range forecast (or, alternatively re-analysis) of meteorological fields and the concomitant prediction of material dispersion based on this forecast or re-analysis (which can result in significant model errors with a complex structure in the predicted concentrations required for the source inversion). Other problems include the fact that the likelihood function in the current source reconstruction problem is an open-form expression whose evaluation is extremely computer intensive. Owing to the fact that simulation-based posterior inference requires that a large number of forward calculations of the source-receptor relationship be undertaken, it is evident that a fast and efficient sampling technique will be needed for the posterior inference. This report describes the methodology for provision of a fast and reliable framework for Bayesian inference in the context of source reconstruction (and, more particularly, in relation to source term estimation problems associated with activity concentration measurements by the radionuclide part of the International Monitoring System that has been set up for verification of the CTBT).

2 Bayesian inference in a nutshell

In a seminal paper, Cox [19] demonstrated that any acceptable calculus of plausible inference must be consistent with a set of three elementary desiderata that conform to the obvious basic properties of inference. This theory can be interpreted as *the* extension of Aristotelian deductive logic to cases where there is uncertainty, and is based on three basic desiderata: namely, (1) degrees of plausibility are represented by real numbers; (2) the measure of plausibility must exhibit a qualitative agreement with rationality (common sense); and, (3) internal consistency in the sense that if a conclusion can be reasoned out in two or more ways, every possible way leads to the same result.

Cox [19] demonstrated that the three desiderata enumerated above are sufficient to determine a quantitative theory of inference in which the rules for manipulating plausibility p reduce simply to the following sum and product rules; namely, the sum rule

$$p(H|I) + p(\overline{H}|I) = 1, \quad (1)$$

and the product rule

$$p(H, D|I) = p(H|I)p(D|H, I) = p(D|I)p(H|D, I). \quad (2)$$

In Eqs. (1) and (2), H , D , and I are arbitrary propositions (hypotheses) and $p(H|I)$ is a measure of the degree to which the hypothesis H is supported by the information embodied in the proposition I , “|” denotes “conditional upon”, $\overline{H}$ denotes “not H ”, and H, D denotes “ H and D ”. Furthermore, it is important to note that the rules embodied in Eqs. (1) and (2) are the ordinary rules of probability calculus. Indeed, Cox demonstrated that every allowed plausibility theory for quantitative inference must be mathematically equivalent (isomorphic) to probability theory, or else inconsistent (in the sense that no other calculus consistent with the above-mentioned desiderata is admissible for inference). A detailed and clear description of the quantitative theory of plausible inference (probability theory as extended logic) has been given by Jaynes [20] in his definitive treatise.

In many cases, the hypothesis H concerns the values of a combination of model parameters denoted by $\boldsymbol{\theta}$ (parameter vector) which we would like to assess in the light of some observed data $D \equiv \mathbf{d}$ (data vector) and any background information I . More specifically, we are interested in calculating the probability of the hypothesis $H \equiv \boldsymbol{\theta}$, given the data $D \equiv \mathbf{d}$ and any prior information I we may have regarding the hypothesis and data. The key in the evaluation of this conditional probability is the product rule of probability calculus [see Eq. (2)] which, in terms of the context described here, can be expressed as follows:

$$p(\boldsymbol{\theta}|I)p(\mathbf{d}|\boldsymbol{\theta}, I) = p(\mathbf{d}|I)p(\boldsymbol{\theta}|\mathbf{d}, I). \quad (3)$$

Equation (3) is the fundamental relationship of Bayesian probability theory (in which probability is interpreted as the degree of belief about a proposition or hypothesis). More specifically, this equation is the expression of Bayes’ theorem which describes a type of learning: how must the probability of a hypothesis concerning $\boldsymbol{\theta}$ be updated in the presence of the new information embodied in the observed data $\mathbf{d}$. Finally, as emphasized by Jaynes [20], in Bayesian probability theory it is the probability that is distributed and not the parameter vector $\boldsymbol{\theta}$. Our incomplete knowledge about $\boldsymbol{\theta}$ is expressed by spreading our belief regarding the true value of the parameter vector among various hypotheses in accordance to the posterior distribution (see below).

The input quantities for Bayesian inference are on the left-hand-side of Eq. (3) and are as follows: $p(\boldsymbol{\theta}|I)$ is the prior probability for a hypothesis about the values of the parameter vector $\boldsymbol{\theta}$ which encodes our state of knowledge about these parameters before the receipt of the data $\mathbf{d}$; and, $p(\mathbf{d}|\boldsymbol{\theta}, I)$ is the likelihood function and is considered to be a function of $\boldsymbol{\theta}$ for fixed data $\mathbf{d}$. The likelihood function incorporates the information provided by the measured data $\mathbf{d}$ into the inferential scheme.

The output quantities provided by the Bayesian inference are on the right-hand side of Eq. (3) and are as follows: $p(\mathbf{d}|I)$ is referred to as the evidence or global likelihood and $p(\boldsymbol{\theta}|\mathbf{d}, I)$ is the posterior probability for the hypothesis about the values of the parameter vector $\boldsymbol{\theta}$, evaluated in light of the additional information provided by the

measured data $\mathbf{d}$. In the context of the determination of the plausible values for the parameter vector (parameter estimation problem), the evidence is simply a normalization constant that is *independent* of the parameter vector and, as a consequence, can therefore be ignored. More specifically, within the context of a parameter estimation problem, the evidence is a constant given by the following multi-dimensional integral over the parameter space [cf. Eq. (3)]:

$$Z \equiv p(\mathbf{d}|I) = \int p(\boldsymbol{\theta}|I)p(\mathbf{d}|\boldsymbol{\theta}, I) d\boldsymbol{\theta}, \quad (4)$$

which ensures the proper normalization of the posterior distribution $p(\boldsymbol{\theta}|\mathbf{d}, I)$ (viz., the condition that the integral of the posterior distribution over its domain of definition is unity). However, it should be stressed that the evidence assumes a central role in the problem of model selection [20] and in this context this quantity cannot be ignored as in the parameter estimation problem considered in this report.

In summary, in Bayesian inference we need to consider both parameters $\boldsymbol{\theta}$ and data $\mathbf{d}$ as well as a model $\mathcal{M}$ that relates the parameters to the data, all within an overarching context I (contextual or background information available in the problem). The joint probability distribution of the parameters $\boldsymbol{\theta}$ and the data $\mathbf{d}$ factorizes, in accordance to the product rule of Eq. (2), in two different ways as summarized by Bayes' rule given in Eq. (3). On the left-hand-side of this rule are the two inputs: the prior distribution for $\boldsymbol{\theta}$ and the likelihood function representing the probability distribution of the data $\mathbf{d}$ for a given value of $\boldsymbol{\theta}$. On the right-hand-side of this rule are the two outputs: the posterior distribution of $\boldsymbol{\theta}$ representing our inference for the parameters after incorporating the information embodied in $\mathbf{d}$ and the evidence which is the probability with which one would predict $\mathbf{d}$ given only prior information about the model $\mathcal{M}$ (implicit in the background information I).

The transformation from the prior distribution to the posterior distribution for $\boldsymbol{\theta}$ is mediated through the factor $p(\mathbf{d}|\boldsymbol{\theta}, I)/p(\mathbf{d}|I)$ (corresponding to a learning process). The quantitative measure of the gain in information content (concerning the parameters $\boldsymbol{\theta}$) obtained from the receipt of the data $\mathbf{d}$ is the information gain (or negative entropy) defined as the “difference” between the prior and the posterior as follows:

$$\mathcal{S} = \int p(\boldsymbol{\theta}|\mathbf{d}, I) \log \left(\frac{p(\boldsymbol{\theta}|\mathbf{d}, I)}{p(\boldsymbol{\theta}|I)} \right) d\boldsymbol{\theta}. \quad (5)$$

The information gain (also known in statistics as the Kullback-Leibler divergence [21]) can be interpreted as the logarithm of the volumetric factor by which the prior has been compressed to become the posterior (the greater this compression, the greater is the information gain provided by the data $\mathbf{d}$). In simpler terms, $\mathcal{S}$ is the amount of information or “surprise” contained in the data $\mathbf{d}$. Note that the update in our state of knowledge from the prior to the posterior distribution (learning process) is

obtained through the modulation of the prior by the ratio of the likelihood function and the evidence ($p(\mathbf{d}|\boldsymbol{\theta}, I)/p(\mathbf{d}|I)$) and that this same ratio is involved in the measure of the information gain $\mathcal{S}$ achieved in this learning process (since $\mathcal{S}$ involves $p(\boldsymbol{\theta}|\mathbf{d}, I)/p(\boldsymbol{\theta}|I) = p(\mathbf{d}|\boldsymbol{\theta}, I)/p(\mathbf{d}|I)$ by Bayes' rule).

3 Assignment of the direct probabilities

In view of the fact that $p(\mathbf{d}|I)$ is merely a normalization constant in the context of the parameter estimation problem, the key output of the Bayesian inference methodology in this case is the posterior probability $p(\boldsymbol{\theta}|\mathbf{d}, I)$ which embodies all the information about the unknown source parameters $\boldsymbol{\theta}$. To determine this quantity requires assigning appropriate functional forms for the two input quantities for Bayesian inference; namely, the prior probability $p(\mathbf{d}|I)$ and the likelihood function $p(\mathbf{d}|\boldsymbol{\theta}, I)$. The assignment of numerical values for these two probabilities in the context of the source estimation problem defines the “vocabulary” (or “what is”) for the rational inference [with the “grammar” (or “how to”) for the rational inference determined by the two rules for combining probabilities embodied by Eqs. (1) and (2)]. In this sense, the two input probabilities in Eq. (3) that need to be assigned directly are the *direct* probabilities for the problem.

3.1 Assignment of likelihood – input

In this report, we focus on the source reconstruction problem: namely, the characterization of the properties of an unknown source following event detection by a network of concentration sensors. The model equation for this problem is given by

$$\begin{aligned} d_J &\equiv d(\mathbf{x}_J, t_J) \\ &= \bar{C}(\boldsymbol{\theta}; \mathbf{x}_J, t_J) + e(\mathbf{x}_J, t_J) \\ &\equiv \bar{C}_J(\boldsymbol{\theta}) + e_J, \end{aligned} \tag{6}$$

$J = 1, 2, \dots, N$ where N is the total number of measured concentration data. In Eq. (6), the concentration data acquired or measured (by a sensor) at the space-time point $(\mathbf{x}_J, t_J)$ is represented by $d(\mathbf{x}_J, t_J)$ and the associated model prediction for the concentration data is given by $\bar{C}(\boldsymbol{\theta}; \mathbf{x}_J, t_J)$. As applied to the problem of source reconstruction, $\boldsymbol{\theta}$ is identified with the set of parameters used to characterize the source distribution (e.g., location, emission rate); $\mathbf{d}$ is associated with the measured concentration data, so $\mathbf{d} \equiv (d_1, d_2, \dots, d_N)$; and, I corresponds to any contextual information that defines the source reconstruction problem (e.g., background meteorology, atmospheric dispersion model $\mathcal{M}$ used to define the source-receptor relationship).

The error (discrepancy) between the measured concentration d_J and the predicted concentration $\bar{C}_J(\boldsymbol{\theta})$ is represented symbolically in Eq. (6) by e_J . In the problem

considered in this paper, the data d_J can differ from the model $\bar{C}_J(\boldsymbol{\theta})$ because of two major contributions to the error e_J : namely, the observation or instrument error in the measured concentration d_J arising from the noise inherent in the sensor and the model error in the determination of the predicted concentration $\bar{C}_J(\boldsymbol{\theta})$. Of these two errors, the model error is by far the most dominant contribution to the total error e_J and is also the most difficult to characterize.

In our current application, the model error arises from three primary sources. These are as follows:

1. uncertainties in the representation of various physical processes in the dispersion model;
2. uncertainties in the input meteorological fields (initial and boundary conditions) used to “drive” the dispersion model (either numerical weather prediction uncertainties if these fields are obtained as a forecast, or data assimilation uncertainties for the state of the atmosphere if these fields are obtained through a re-analysis); and,
3. uncertainties in the numerical solution of the model equations that characterize the dispersion model (which include both discretization errors and statistical model errors, the latter of which arises from using necessarily a finite number of “marked” fluid particles to estimate the mean concentration field in the case of a Lagrangian stochastic model of dispersion).

As a consequence of the complexity in structure of the error e_J [cf. Eq. (6)] for our current application, it is extremely difficult (if not insuperable) to specify *a priori* an exact value σ_J for the standard deviation of e_J ($J = 1, 2, \dots, N$). If the standard deviation σ_J of e_J was exactly known, then it can be shown by application of the principle of maximum entropy that a Gaussian distribution of the form

$$p(\mathbf{d}|\boldsymbol{\theta}, I) = \frac{1}{\prod_{J=1}^N \sqrt{2\pi}\sigma_J} \exp\left(-\frac{1}{2}\chi^2(\boldsymbol{\theta})\right) \quad (7)$$

with

$$\chi^2(\boldsymbol{\theta}) \equiv \sum_{J=1}^N \left(\frac{d_J - \bar{C}_J(\boldsymbol{\theta})}{\sigma_J} \right)^2 \quad (8)$$

would be the most conservative choice for the direct probability (or likelihood) of the concentration data $\mathbf{d}$ (see Jaynes [20]). Indeed, assigning a Gaussian distribution for the noise e_J using the principle of maximum entropy makes no statement about the true (unknown) sampling distribution of the noise. Rather, it simply represents a maximally uninformative state of knowledge, a state of knowledge that reflects what the observer knows about the true noise in the data (namely, the mean and variance of the noise, with all other properties of the noise being irrelevant to the inference

since these are unknown to the observer). Finally, it should be emphasized that the standard deviations $\{\sigma_J, J = 1, 2, \dots, N\}$ (which are taken in Eqs. (7) and (8) to be known) are not assumed to be equal for each data point (modeled and measured), so the model is heteroscedastic (meaning that the squared errors $(d_J - \bar{C}_J(\boldsymbol{\theta}))^2$ have different expected values σ_J^2).

However, in our current application, we do not know the true or actual standard deviation (or, equivalently variance) of the noise. Because σ_J is not known *a priori*, it is useful to characterize the uncertainty in the specification of σ_J with a probability distribution. Following Yee [16], we will choose this probability distribution to be an inverse gamma distribution with the following form:

$$\varphi(\sigma_J | s_J, \alpha, \beta) = 2 \frac{\alpha^\beta}{\Gamma(\beta)} \left(\frac{s_J}{\sigma_J} \right)^{2\beta} \exp \left(-\alpha \frac{s_J^2}{\sigma_J^2} \right) \frac{1}{\sigma_J}, \quad J = 1, 2, \dots, N. \quad (9)$$

Here, $\Gamma(x)$ denotes the gamma function, α and β are scale and shape parameters of the inverse gamma distribution, and s_J is the quoted (nominal) estimate for the true but unknown standard deviation σ_J . The values for the hyperparameters α and β are chosen as $\alpha = \pi^{-1}$ and $\beta = 1$ following the rationale described in Yee [16].

The likelihood function in Eqs. (7) and (8) depends on the error standard deviations σ_J ($J = 1, 2, \dots, N$) which are generally unknown. To remove these unwanted parameters (nuisance parameters), we can multiply the likelihood given in Eq. (7) by the (assigned) probability distribution for each of the error standard deviations embodied in Eq. (9) and integrate the result with respect to the unwanted parameters (error standard deviations) to give an integrated likelihood function with the following form [16]:

$$\begin{aligned} p(\mathbf{d} | \boldsymbol{\theta}, \mathbf{s}, \alpha, \beta) &= \int p(\mathbf{d} | \boldsymbol{\theta}, I) \prod_{J=1}^N \varphi(\sigma_J | s_J, \alpha, \beta) d\boldsymbol{\sigma} \\ &= \prod_{J=1}^N \frac{\alpha^\beta \Gamma(\beta + 1/2)}{\sqrt{2\pi} s_J \Gamma(\beta)} \left[\alpha + (d_J - \bar{C}_J(\boldsymbol{\theta}))^2 / (2s_J^2) \right]^{-\beta-1/2}. \end{aligned} \quad (10)$$

In Eq. (10), $\mathbf{s} \equiv (s_1, s_2, \dots, s_N)$ is the vector of estimated (quoted) standard deviations for the error e_J ($J = 1, 2, \dots, N$) and $d\boldsymbol{\sigma} \equiv d\sigma_1 d\sigma_2 \dots d\sigma_N$. The process of integrating out the nuisance parameters in Eq. (10) is called marginalization, which provides the general recipe for the elimination of unwanted nuisance variables from a Bayesian calculation. This represents simply the application of the sum rule of probability theory.

To be more specific with respect to the source parameter vector $\boldsymbol{\theta}$, a source distribution S with the following general form is considered in this paper:

$$S(\mathbf{x}, t) = \sum_{k=1}^{N_s} Q_{s,k} \mathbb{I}_{[t_{k-1}, t_k)}(t) \delta(\mathbf{x} - \mathbf{x}_s), \quad (11)$$

where $\delta(x)$ is the Dirac delta function and $\mathbb{I}_A(x)$ denotes the indicator function for set A defined as $\mathbb{I}_A(x) = 1$ if $x \in A$ and $\mathbb{I}_A(x) = 0$ if $x \notin A$. In Eq. (11), $Q_{s,k}$ is the (constant) source emission rate over the time interval $[t_{k-1}, t_k)$ ($k = 1, 2, \dots, N_s$) with $t_i < t_j$ for $i < j$. It is assumed that temporal change-point positions $\{t_j, j = 0, 1, \dots, N_s\}$ that define the constant emission rate sections are fixed and known. Equation (11) describes a (continuous) point source located at the vector position $\mathbf{x}_s$ whose emission rate temporal profile is described by a piecewise constant function characterized by $\{Q_{s,k}, k = 1, 2, \dots, N_s\}$. The parameters describing this source distribution can be assembled into a source parameter vector given by $\boldsymbol{\theta} \equiv (\mathbf{x}_s, Q_{s,1}, Q_{s,2}, \dots, Q_{s,N_s}) \in \mathbb{R}^{3+N_s}$.

3.2 Assignment of prior – input

To assign the prior probability for the source parameters $\boldsymbol{\theta}$, it is necessary to state explicitly what is known about these parameters. Firstly, it is assumed that the source parameters are logically independent with the result that the prior distribution $p(\boldsymbol{\theta}|I)$ is given by the product of the individual prior distributions:

$$p(\boldsymbol{\theta}|I) = p(\mathbf{x}_s|I) \prod_{k=1}^{N_s} p(Q_{s,k}|I). \quad (12)$$

Secondly, the source location and emission rates are known *a priori* to be bounded. In particular, it is assumed that the location $\mathbf{x}_s$ of the source is contained in some spatial region $\mathcal{D} \subset \mathbb{R}^3$. Furthermore, the emission rate $Q_{s,k}$ is assumed to be bounded by $Q_{\min} < Q_{s,k} < Q_{\max}$ ($k = 1, 2, \dots, N_s$) where $Q_{\min}$ and $Q_{\max}$ are the lower and upper bounds for each emission rate, respectively. Different lower and upper bounds can be chosen for each emission rate $Q_{s,k}$ in the assignment of the prior distribution for this quantity, but for the formulation herein we simply use a common lower and upper bound for all emission rates (with effectively no loss in generality).

Finally, if nothing else is known about the various source parameters except for their lower and upper bounds, then application of the principle of maximum entropy to our state of knowledge concerning the source parameters, results in the assignment of a uniform prior distribution for these parameters, so

$$p(\boldsymbol{\theta}|I) \propto \mathbb{I}_{\mathcal{D}}(\mathbf{x}_s) \prod_{k=1}^{N_s} \mathbb{I}_{(Q_{\min}, Q_{\max})}(Q_{s,k}). \quad (13)$$

3.3 Posterior distribution – output

Combining Eqs. (10) and (13) with reference to Bayes' rule given by Eq. (3), yields the key output of the Bayesian inference; namely, the posterior distribution $p(\boldsymbol{\theta}|\mathbf{d}, I)$

which for our current application assumes the following form:

$$p(\boldsymbol{\theta}|\mathbf{d}, \mathbf{s}, \alpha, \beta, I) \propto \mathbb{I}_{\mathcal{D}}(\mathbf{x}_s) \prod_{k=1}^{N_s} \mathbb{I}_{(Q_{\min}, Q_{\max})}(Q_{s,k}) \times \prod_{J=1}^N \frac{\alpha^\beta \Gamma(\beta + 1/2)}{\sqrt{2\pi} s_J \Gamma(\beta)} \left[\alpha + (d_J - \bar{C}_J(\boldsymbol{\theta}))^2 / (2s_J^2) \right]^{-\beta-1/2}. \quad (14)$$

Note that the quantities $\mathbf{s}$, α and β have been added explicitly to the posterior probability of the source parameters $\boldsymbol{\theta}$ in Eq. (14) to indicate that these quantities are known (viz., they are provided *a priori* by the user, in addition to the measured concentration data $\mathbf{d}$).

The posterior distribution given by Eq. (14) provides the full solution for the source reconstruction (or source term estimation) problem. Inferences on the values of the source parameters are based on this posterior distribution. Indeed, it is useful to summarize the posterior distribution in terms of a few relevant quantities. One relevant quantity is the posterior mean of the various source parameters (e.g., location, emission rates) given by

$$\langle \theta_j \rangle = \mathbb{E}[\theta_j | \mathbf{d}, \mathbf{s}, \alpha, \beta, I] \equiv \int \theta_j p(\boldsymbol{\theta} | \mathbf{d}, \mathbf{s}, \alpha, \beta, I) d\boldsymbol{\theta}, \quad (15)$$

where θ_j is the j -th component of $\boldsymbol{\theta}$ and $\mathbb{E}[\cdot]$ denotes mathematical expectation. A measure of the uncertainty of this estimate (posterior mean) for θ_j is given the posterior standard deviation

$$\begin{aligned} \sigma(\theta_j) &= \left(\mathbb{E}[(\theta_j - \langle \theta_j \rangle)^2 | \mathbf{d}, \mathbf{s}, \alpha, \beta, I] \right)^{1/2} \\ &= \left(\int (\theta_j - \langle \theta_j \rangle)^2 p(\boldsymbol{\theta} | \mathbf{d}, \mathbf{s}, \alpha, \beta, I) d\boldsymbol{\theta} \right)^{1/2}. \end{aligned} \quad (16)$$

Alternatively, a $100r\%$ [$r \in (0, 1)$] credible or highest posterior distribution (HPD) region for $\boldsymbol{\theta}$ can be used as a measure of the uncertainty in the determination of the source parameter vector. The $100r\%$ HPD region is defined in terms of that portion $\mathcal{R}$ of the source parameter space such that

$$\int_{\mathcal{R}} p(\boldsymbol{\theta} | \mathbf{d}, \mathbf{s}, \alpha, \beta, I) d\boldsymbol{\theta} = r, \quad (17)$$

with the posterior density within $\mathcal{R}$ everywhere larger than outside it. This enables one to state that the probability that $\boldsymbol{\theta}$ lies within $\mathcal{R}$, given the observed data $\mathbf{d}$ and background information I , is at least $100r\%$.

4 Bayesian computation

There are two computational problems that need to be addressed before the Bayesian inferential methodology for source reconstruction can be applied to practical real-world applications; namely, (1) Bayesian inversion of concentration data requires a fast and efficient technique for the determination of the source-receptor relationship (viz., for the rapid computation of $\bar{C}_J(\boldsymbol{\theta})$ for a given hypothesis about the source parameters $\boldsymbol{\theta}$); and, (2) methodology for the efficient sampling of the posterior distribution $p(\boldsymbol{\theta}|\mathbf{d}, \mathbf{s}, \alpha, \beta, I)$. We will address each of these computational problems in turn.

4.1 Rapid computation of model concentration

The Bayesian inferential methodology for source term estimation is highly computer intensive because the simulation-based inference procedure requires a large number of calculations of the mean concentration $\bar{C}_J(\boldsymbol{\theta})$ ($J = 1, 2, \dots, N$) to be determined. These calculations for the mean concentration need to be executed for a large number of source parameter hypotheses $\boldsymbol{\theta}$ required for the complete exploration of the posterior distribution $p(\boldsymbol{\theta}|\mathbf{d}, \mathbf{s}, \alpha, \beta, I)$ (see below). As a consequence, for the Bayesian inversion of concentration data to be practical, fast and efficient techniques are required for the determination of the source-receptor relationship within the context of sampling in the hypothesis space of the source parameters required for posterior inference.

For the current application, it is possible to construct an emulator for the (concentration) simulation model and use this emulator as a computationally inexpensive surrogate to replace the forward (source-oriented) atmospheric dispersion model used normally to determine the source-receptor relationship [22]. However, applying this type of approximation is not required for the current problem. It was shown by Keats et al. [7] and Yee et al. [11] that an *exact* computationally efficient procedure (appropriate for use in a Bayesian inference scheme) exists in the form of a receptor-oriented scheme for the representation of the source-receptor relationship.

Using a standard forward atmospheric dispersion model (say, a forward Lagrangian stochastic (LS) model) the predicted concentration $\bar{C}_J(\boldsymbol{\theta})$ “seen” by a sensor can be computed by averaging the mean concentration $C(\mathbf{x}, t)$ obtained from the forward model over the sensor volume and sampling time to give

$$\begin{aligned}\bar{C}_J(\boldsymbol{\theta}) &\equiv \bar{C}(\boldsymbol{\theta}; \mathbf{x}_J, t_J) \\ &= \int_{-\infty}^{t_J} dt \int_{\mathcal{D}} d\mathbf{x} C(\mathbf{x}, t) h(\mathbf{x}, t | \mathbf{x}_J, t_J) \equiv \langle C|h \rangle(\mathbf{x}_J, t_J),\end{aligned}\quad (18)$$

where $h(\mathbf{x}, t | \mathbf{x}_J, t_J)$ is the spatial-temporal filtering function for the sensor concentra-

tion measurement at $(\mathbf{x}_J, t_J)$ and $\mathcal{D} \times (-\infty, t_J)$ corresponds to the space-time volume that contains the source distribution and the sensors (receptors). Note that the concentration $\bar{C}_J(\boldsymbol{\theta})$ “seen” by a sensor can be expressed as the inner product $\langle C|h \rangle$ of the mean concentration C and the sensor response function h . Alternatively, it should be noted that with the source distribution S given by Eq. (11), the predicted concentration at a receptor location $\mathbf{x}_J$ and time t_J can also be determined from the following dual relationship once the adjunct (or dual) concentration field C^* has been computed for this location and time:

$$\begin{aligned}\bar{C}_J(\boldsymbol{\theta}) &= \int_{-\infty}^{t_J} dt' \int_{\mathcal{D}} d\mathbf{x}' C^*(\mathbf{x}', t' | \mathbf{x}_J, t_J) S(\mathbf{x}', t') \\ &\equiv \langle C^* | S \rangle(\mathbf{x}_J, t_J) = \langle C | h \rangle(\mathbf{x}_J, t_J),\end{aligned}\quad (19)$$

where $C^*(\mathbf{x}', t' | \mathbf{x}_J, t_J)$ is the adjunct concentration at space-time point $(\mathbf{x}', t')$ associated with the sensor concentration datum measured at location $\mathbf{x}_J$ and time t_J .

It is noted that the predicted mean concentration $\bar{C}_J(\boldsymbol{\theta})$ “seen” by a sensor for a given hypothesis about the source parameters $\boldsymbol{\theta}$ can be rapidly computed by simply evaluating the inner (or scalar) product $\langle C^* | S \rangle$ of the adjunct concentration C^* and the source distribution S corresponding to the given hypothesis. In other words, the predicted concentration $\bar{C}_J$ is obtained from a mathematical model (say, a backward Lagrangian stochastic model for dispersion) by evaluation of the bounded linear functionals $\bar{C}_J = \langle C_J^* | S \rangle$ for $J = 1, 2, \dots, N$ where C_J^* denotes the adjunct concentration field obtained at the sensor space-time point $(\mathbf{x}_J, t_J)$ and where S corresponds to a given hypothesis about the source. In applied mathematics, the elements C_J^* are referred to usually as representers. More specifically, if we substitute Eq. (11) into Eq. (19), the model (predicted) concentration $\bar{C}_J(\boldsymbol{\theta})$ “seen” by the sensor at location $\mathbf{x}_J$ and time t_J is given explicitly by

$$\bar{C}_J(\boldsymbol{\theta}) = \sum_{k=1}^{N_s} Q_{s,k} \int_{t_{k-1}}^{\min(t_J, t_k)} C^*(\mathbf{x}_s, t' | \mathbf{x}_J, t_J) dt'. \quad (20)$$

For the application reported herein, a backward Lagrangian stochastic model for long-range transport was used to determine the adjunct concentration field C^* over the northern hemisphere. The backward LS model employed here is an operational model used by the Canadian Meteorological Centre to support both Canadian Treaty monitoring and the various mandates of the Provisional Technical Secretariat of the Comprehensive Nuclear-Test-Ban Treaty Organization, including event analysis by member states. The backward LS model for the determination of C^* was “driven” by re-analyzed meteorological fields that were obtained at a relatively low temporal and spatial resolution; namely, at a temporal resolution of 6 h (rate of the data assimilation) with a core spatial resolution of 0.5° on a geographical latitude and longitude coordinate system. The backward LS model was run retrospectively using

these re-analyzed meteorological fields, providing C^* fields with a temporal resolution of 3 h over a period of 14 d prior to the commencement of the sampling for a particular activity concentration measurement.

4.2 Markov chain Monte Carlo

An examination of Eq. (14) shows that the posterior distribution is highly nonlinear in the source parameters and explicit evaluations of Eqs. (15), (16) and (17) (and similar multi-dimensional integrals required for a full Bayesian computation of desired marginal posteriors and various Bayesian statistics) are impossible. In view of this insuperable difficulty, we apply posterior sampling for the evaluation of these integrals which is implemented using a Markov chain Monte Carlo (MCMC) algorithm (see Gilks et al. [23] and Gelman et al. [24]). A MCMC algorithm will be used to generate samples of source distribution models (characterized by θ) from the posterior distribution given by Eq. (14). In other words, the objective here is to find an ensemble of source distribution models that preferentially sample the high plausibility regions of the source parameter space (as determined by the posterior distribution of θ), rather than seek a single optimal model. Towards this objective, a Markov chain $\{\theta^{(t)}\}$ is constructed whose stationary (or invariant) distribution is the posterior distribution $p(\theta|\mathbf{d}, \mathbf{s}, \alpha, \beta, I)$ of the parameters θ .

All quantities of interest, such as the posterior means and standard deviations of the various source parameters and the various marginal posterior distributions, can be estimated by sample path averages of the Markov process $\{\theta^{(t)}\}$. In general, MCMC algorithms are used to generate the required Markov process $\{\theta^{(t)}\}$. The Metropolis-Hastings (M-H) algorithm [23] forms the underlying basis for MCMC sampling and it is perhaps not too surprising that the M-H algorithm has become almost synonymous with MCMC sampling. Indeed, most of the algorithms for MCMC sampling reported in the literature [23, 24] can be interpreted as either special cases or extensions of the basic M-H algorithm. The M-H algorithm generates the required samples (ensemble of source models) by constructing a kind of random walk in the source parameter space so that the probability of being in a particular region of this space is proportional to the posterior probability mass for that region.

The basic M-H algorithm for generating this random walk consists of two components: (1) a proposal distribution $q(\theta'|\theta)$; and, (2) an acceptance probability $\alpha(\theta, \theta')$. These two components determine the two steps of the algorithm. Firstly, given a chain in the current state $\theta^{(t)} = \theta$ at time (iteration) t , a proposed new state θ' is drawn from the proposal distribution $q(\theta'|\theta)$. Secondly, this new point is accepted or rejected as the new state of the chain at time $(t + 1)$ using the standard M-H acceptance

probability [23] given by

$$\alpha(\boldsymbol{\theta}, \boldsymbol{\theta}') = \min \left\{ 1, \frac{p(\boldsymbol{\theta}'|\mathbf{d}, \mathbf{s}, \alpha, \beta, I)q(\boldsymbol{\theta}|\boldsymbol{\theta}')}{p(\boldsymbol{\theta}|\mathbf{d}, \mathbf{s}, \alpha, \beta, I)q(\boldsymbol{\theta}'|\boldsymbol{\theta})} \right\}. \quad (21)$$

If the proposal is accepted, then $\boldsymbol{\theta}^{(t+1)} = \boldsymbol{\theta}'$; otherwise, $\boldsymbol{\theta}^{(t+1)} = \boldsymbol{\theta}$. The construction of the M-H algorithm ensures that the condition of detailed balance holds, implying that the M-H algorithm converges to a stationary distribution which corresponds to the target distribution ($p(\boldsymbol{\theta}|\mathbf{d}, \mathbf{s}, \alpha, \beta, I)$) that we are trying to draw samples from.

The simple M-H algorithm summarized above is subject to various difficulties if the target probability distribution is multi-modal or possesses significant curving degeneracies in the (possibly) high-dimensional parameter space. To overcome these difficulties, a number of approaches has been proposed to increase the sampling efficiency of MCMC simulation based on the M-H algorithm. In the current study, two different approaches for improving the posterior exploration of source distribution models (allowing the parameter space to be explored more freely than in the standard M-H algorithm) have been used.

4.2.1 Differential evolution adaptive Metropolis sampling

One of the MCMC algorithms used for the current study is a multiple-try differential evolution adaptive Metropolis algorithm with sampling from an archive of past states (MT-DREAM_(ZS)). The details of this MCMC sampling algorithm are described by Laloy and Vrugt [25], but for completeness we will briefly summarize the main components of this algorithm. In particular, only the relevant details of the algorithm that are required for the interpretation of the results in this paper are emphasized.

Firstly, MT-DREAM_(ZS) samples from an archive of past states $\mathbf{Z}$ to generate the candidate points (proposals) for each of the individual Markov chains that are used to explore the target posterior distribution. The archive of past states is initialized by drawing M_0 samples of $\boldsymbol{\theta}$ from the prior distribution $p(\boldsymbol{\theta}|I)$ [see Eq. (13)]. The states of the individual Markov chains at various times (iterations) are periodically stored in the archive $\mathbf{Z}$ using a simple thinning rule (e.g., after every $K \geq 1$ iterations, the states from the individual Markov chains are added to the archive $\mathbf{Z}$).

Secondly, as already alluded to here, the algorithm utilizes multiple (different) Markov chains that are run simultaneously in parallel. These multiple chains employ a self-adaptive randomized subspace sampling of difference vectors from the archive $\mathbf{Z}$ of past states to generate new candidate points in these chains. More specifically, if Markov chain i is in the current state $\boldsymbol{\theta}_i^{(t)} \equiv \boldsymbol{\theta}_i$, a proposed candidate state $\boldsymbol{\theta}'_i$ is

generated in accordance to the following prescription:

$$\boldsymbol{\theta}'_i = \boldsymbol{\theta}_i + (\mathbf{1}_{N_p} + \mathbf{e}_{N_p})\gamma(\delta, N'_p) \left[\sum_{j=1}^{\delta} \mathbf{Z}_{r_1(j)} - \sum_{n=1}^{\delta} \mathbf{Z}_{r_2(n)} \right] + \boldsymbol{\epsilon}_{N_p}, \quad (22)$$

for $i = 1, 2, \dots, N_c$. Here, N_c and N_p are the number of Markov chains and the dimensionality of the parameter vector $\boldsymbol{\theta}_i$, respectively; $\mathbf{Z}_k$ ($k \in \mathbb{N}$) denotes the k -th state vector stored in the archive $\mathbf{Z}$; $r_1(j)$ and $r_2(n)$ are random integers drawn from the set $\{1, 2, \dots, M\}$ where M is the number of samples in the archive with $r_1(j) \neq r_2(n)$ ($j, n = 1, 2, \dots, \delta$ and δ is the dimension of the randomized subspace); $\mathbf{1}_{N_p}$ is a vector of dimension N_p consisting of ones; $\mathbf{e}_{N_p}$ and $\boldsymbol{\epsilon}_{N_p}$ are random vectors of dimension N_p whose components are drawn from $U(-b, b)$ (uniform distribution) and $N(0, b^*)$ (standard normal distribution), respectively, where b and b^* are small compared to the “width” of the target distribution; and, γ is the size of the update step whose value depends on δ and N'_p (N'_p is the number of dimensions of $\boldsymbol{\theta}_i$ that is updated using the binomial scheme described below). An appropriate choice for γ is given in [25].

Thirdly, as part of the randomized subspace sampling strategy, each element of the candidate points for the parallel proposals is updated (accepted) in accordance to a binomial scheme (Bernoulli trial) with a crossover probability p_s , otherwise the proposed element retains its previous (old) value. This multiple-chain approach automatically adjusts or adapts the scale and orientation of the proposal function. In addition to the randomized subspace update given in Eq. (22), the MT-DREAM_(ZS) algorithm also employs a snooker step update of the state with an adaptive step size. This update is included with a fixed (albeit small) probability in order to improve the mixing efficiency of the algorithm for exploration of the hypothesis space.

Fourthly, to further improve the efficiency of the sampling, MT-DREAM_(ZS) incorporates a multiple-try Metropolis (MTM) approach proposed initially by Liu et al. [26]. The basic idea underpinning the MTM approach is as follows: longer range candidate moves are rarely accepted, but if multiple points are proposed for these longer range moves then the acceptance probability will be increased. The MTM algorithm is applied individually to each of the different Markov chains used in the MT-DREAM_(ZS), involving generating l draws using the randomized subspace sampling procedure for each chain [see Eq. (22)], choosing one of these draws (proposals) as the reference point, and generating a new set of $(K - 1)$ draws with respect to this reference point using the randomized subspace sampling strategy. The acceptance rule is simply the Metropolis-Hasting acceptance probability [see Eq. (21)] applied to the sequence of proposals that comprise the MTM schema.

Finally, to determine if the multiple Markov chains used in MT-DREAM_(ZS) have achieved stationarity (viz., have converged to the stationary distribution associated

with the chains), the Gelman and Rubin [27] convergence diagnostic $\hat{R}$ is computed for each dimension of each chain using the last 50% of the samples of the chain. This simple and popular convergence diagnostic determines whether a chain has achieved stationarity by comparing the variance within each chain and the variance between chains (interchain variance).

4.2.2 Nested sampling

Nested sampling is an innovative Monte Carlo methodology developed by Skilling [28] for general Bayesian computation. The nested sampling algorithm focuses on the computation of the evidence integral [see Eq. (4)]. The basic idea is to transform this multi-dimensional integral over the source parameter space into a simple one-dimensional integral given by

$$Z \equiv p(\mathbf{d}|I) = \int_0^1 \mathcal{L}(\chi) d\chi, \quad (23)$$

where $\mathcal{L}(\chi^*)$ is the value of the likelihood function such that the volume in the parameter space enclosed by the prior $p(\boldsymbol{\theta}|I)$ satisfying $p(\mathbf{d}|\boldsymbol{\theta}, I) \geq \mathcal{L}(\chi^*) \equiv \mathcal{L}^*$ is χ^* . More specifically,

$$\chi^* = \chi(\mathcal{L}^*) = \int_{p(\mathbf{d}|\boldsymbol{\theta}, I) > \mathcal{L}^*} p(\boldsymbol{\theta}|I) d\boldsymbol{\theta}, \quad (24)$$

is the prior mass in the parameter (hypothesis) space enclosed with a likelihood greater than $\mathcal{L}^*$ and $\mathcal{L}(\chi)$ is the inverse which labels the likelihood contour that encloses a prior mass χ . In other words, we label each element of prior mass $d\chi(\boldsymbol{\theta}) = p(\boldsymbol{\theta}|I)d\boldsymbol{\theta}$, and Eq. (23) represents only a change in notation for the evidence integral.

With the change of variables used in Equation (23), the evidence integral is a one-dimensional integral which is conceptually easy to approximate numerically. In particular, if we can evaluate the likelihood function values $\mathcal{L}(\chi)$ at a sequence of m points χ_i ($i = 1, 2, \dots, m$) ordered such that $0 < \chi_m < \chi_{m-1} < \dots < \chi_2 < \chi_1 < 1$ with $\mathcal{L}(\chi_i) > \mathcal{L}(\chi_j)$ ($i > j$) [since $\mathcal{L}$ is necessarily a monotonically decreasing function of χ], the evidence Z can be approximated from the likelihood-ordered samples using the following quadrature rule:

$$\hat{Z} = \sum_{i=1}^m \mathcal{L}(\chi_i) w_i, \quad w_i = (\chi_{i-1} - \chi_i), \quad (25)$$

where the quadrature weights given here correspond to the simple (naïve) rectangle rule. Indeed, a perusal of Eqs. (23) and (24) reveals that χ can be interpreted as the cumulative distribution function of the prior distribution corresponding to a particular ordering of values of the likelihood function.

The nested sampling algorithm provides a set of points $\{\boldsymbol{\theta}_i, \mathcal{L}_i\}$ ($\mathcal{L}_i \equiv \mathcal{L}(\chi_i)$) and a probability distribution over the corresponding set of prior masses $\{\chi_i\}$ associated

with these points. If the prior mass points χ_i are sampled in a logarithmic manner as $\chi_i = \prod_{j=1}^i t_j$ where $t_j \in (0, 1)$ is a shrinkage ratio, then the nested sampling algorithm consists of the following steps. The reader is referred to Skilling [28] for further details of the algorithm:

1. Initialization: $i = 0$, $\chi_0 = 1$, $Z_0 = 0$, and $f = 0.5$;
2. Draw M samples (“live” points) $\boldsymbol{\theta}$ from prior distribution $p(\boldsymbol{\theta}|I)$;
3. Evaluate likelihood function $p(\mathbf{d}|\boldsymbol{\theta}, I)$ for each of these “live” points;
4. $i \rightarrow i + 1$;
5. Select sample with lowest likelihood (labeled as $\mathcal{L}_i$) and remove (discard) it;
6. Shrink the prior mass as follows: $\chi_i = \chi_{i-1} \exp(-1/M)$;
7. Draw a new sample $\boldsymbol{\theta}$ from prior $p(\boldsymbol{\theta}|I)$ subject to $p(\mathbf{d}|\boldsymbol{\theta}, I) > \mathcal{L}_i$ and add this sample to the “live” points;
8. Increment the evidence: $Z_i = Z_{i-1} + \mathcal{L}_i(\chi_{i-1} - \chi_i)$;
9. Convergence test: if $\mathcal{L}_{\max}\chi_i < fZ_i$, go to Step 10; else, go to Step 4;
10. Evaluate contribution to Z_i from remaining M “live” points; stop.

In Step 1 (initialization), f is a preset fraction that is used as the stopping criterion for the nested sampling algorithm. In Step 9 (convergence test), $\mathcal{L}_{\max}$ is the largest value of the likelihood in the set of “live” points. Note in Step 5 that the first sample discarded $\boldsymbol{\theta}^{(1)}$ corresponds to that sample with the smallest value of the likelihood function from the initial M draws, associated as such to the sample with the largest value of χ . For this parameter to be at χ_1 , the remaining $(M - 1)$ sample points must have $\chi < \chi_1$ implying that $p(\chi_1|M) = M(\chi_1)^{M-1}$. This is a standard result from order statistics [29]. It can be shown [28] that the joint distribution over the entire sequence of prior masses sampled using the recursive procedure of the nested sampling algorithm is given by

$$p(\boldsymbol{\chi}) = M^m (\chi_1)^{M-1} \prod_{i=2}^m \frac{(\chi_i)^{M-1}}{(\chi_{i-1})^M}, \quad (26)$$

where $\boldsymbol{\chi} \equiv (\chi_1, \chi_2, \dots, \chi_m)$.

In Step 10 of the algorithm, when the stopping condition is satisfied, the estimate for the evidence is completed by adding the contribution of the remaining (active or “live”) M samples in the ensemble to Z_i ; namely, $\hat{Z} \rightarrow Z_i + M^{-1}(p(\mathbf{d}|\boldsymbol{\theta}_1, I) + p(\mathbf{d}|\boldsymbol{\theta}_2, I) + \dots + p(\mathbf{d}|\boldsymbol{\theta}_M, I))\chi_i$, where $p(\mathbf{d}|\boldsymbol{\theta}_k, I)$ ($k = 1, 2, \dots, M$) are the likelihood

values evaluated at the remaining M samples. Although the focus of the nested sampling algorithm is to evaluate the evidence Z , it is important to note that the samples discarded in Step 5 of the algorithm are actually weighted samples θ_i drawn from the posterior distribution $p(\theta|\mathbf{d}, I) \propto p(\theta|I)p(\mathbf{d}|\theta, I)$. More specifically, the i th discarded sample θ_i in Step 5 of the algorithm can be interpreted as a sample drawn from the posterior distribution of θ with weight given by $\omega_i = \mathcal{L}_i(\chi_{i-1} - \chi_i)/\hat{Z}$ where $\hat{Z}$ is the estimate of the evidence obtained in Step 10 on termination of the algorithm. Finally, it should be noted that the samples (discarded and active) obtained in the nested sampling algorithm can be used to provide also an estimate for the information gain $\mathcal{S}$ given in Eq. (5), using a quadrature rule similar to that used to approximate the evidence Z .

The nested sampling algorithm assumes in Step 7 that drawing a sample from the prior distribution $p(\theta|I)$ lying within a prescribed hard likelihood constraint given by $p(\mathbf{d}|\theta, I) > \mathcal{L}^*$ is possible. To this purpose, Feroz et al. [30] developed a very efficient algorithm (which the authors refer to as MultiNest) for sampling from a prior within a hard likelihood constraint using a very sophisticated procedure for decomposition of the support of the likelihood above a given bound $\mathcal{L}^*$ into a set of overlapping ellipsoids and then sampling from the resulting ellipsoids. The algorithm is appropriate for sampling from posterior distributions with multiple modes and with pronounced curving degeneracies in a high-dimensional parameter space.

5 Applications

The International Monitoring System (IMS) consists of a comprehensive network of seismic, hydroacoustic, infrasound, and radionuclide sensors as part of the verification regime of the Comprehensive Nuclear-Test-Ban Treaty which bans nuclear explosions. A subset of the IMS is the subnetwork of radionuclide gamma detectors/particle filters for the measurement of the activity concentration for various radionuclides (e.g., particulate or aerosol species such as cesium-137 and iodine-131 and/or noble gases such as xenon-133). The IMS radionuclide network will (eventually) have 80 monitoring stations worldwide for the measurement of the activity concentration for particulate/aerosol radioactive species, of which at least 40 of those stations would also have the capability to measure the activity concentration of noble gases [17]. The stations provide 12 or 24 h-averaged activity concentrations of the various radionuclides depending on the technology used.

In this section, we apply the Bayesian inferential methodology for source reconstruction to two different cases involving concentration measurements made by sensors that form part of the IMS radionuclide network. The first case involves emission rate profile reconstruction (temporal history of the emission rate at a fixed *known* location) using some simulated concentration data and the second case involves source

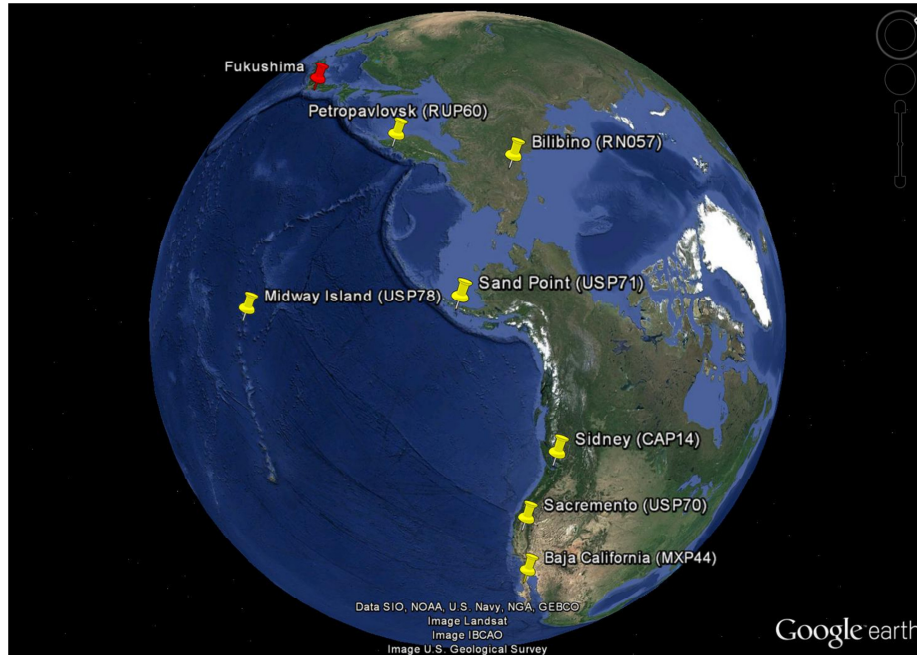


Figure 1: Locations of the seven sampling stations (yellow markers) from the International Monitoring System radionuclide network used for the emission rate profile reconstruction. The location of the Xe-133 tracer source was assumed to be known *a priori* to be at the Fukushima Daiichi nuclear power plant (red marker).

reconstruction (location and emission rate) using some real concentration data.

5.1 Case 1: Emission rate profile reconstruction

In this case, simulated concentration data were generated at seven sampling stations that formed part of the IMS radionuclide network and the Bayesian inference scheme described above was applied to reconstruct the unknown emission rate profile for the release. The contamination events “observed” by these sampling stations have been simulated based on a real-world release; namely, the dispersion of radionuclides and noble gases released during the Fukushima Daiichi nuclear power plant accident.

On 11 March 2011, a magnitude 9.0 earthquake followed by a large tsunami occurred off the Pacific coast of Japan. The tsunami caused the Fukushima Daiichi nuclear power plant to lose electric power, which resulted in the accidental release of a large amount of radioactive material from the plant in the subsequent hours and days after the incident [31]. Initial estimates of the release of the radionuclide materials from the plant were based on the known inventory of the radionuclides in the nuclear reactor and the possible behavior of these materials when subjected to the reactor-core meltdown conditions.

Subsequently, a number of researchers have tried to estimate the source term for this release by combining information from atmospheric dispersion calculations with available environmental monitoring data (which include both air concentration and surface deposition measurements). For the purposes of validating our proposed Bayesian methodology for emission profile reconstruction, we will use the emission rate profile for xenon-133 (Xe-133) from the Fukushima nuclear power plant, as determined by the analysis conducted by Stohl et al. [32], as the “true” (albeit unknown for our purposes) source term. The choice of the emission rate profile of Stohl et al. was for convenience only, and not because it was considered superior to emission rate profiles derived by other research groups.

Owing to the fact that the re-analyzed meteorological fields were limited *only* to a 10-day temporal duration after the earthquake that caused the Fukushima Daiichi nuclear power plant failure, the seven nearest IMS monitoring sites to this power plant were used to synthesize the concentration data for the first case study. It should be noted that in reality not all of these sites have noble gas monitoring capability, but for the purposes of this hypothetical example which was used to provide an initial test of the source reconstruction methodology proposed herein, this should not matter. In consequence, for this (hypothetical) case study, we synthesized artificial concentration data for these seven sampling IMS stations using the “true” source term for the radioxenon release from the Fukushima nuclear power plant. These seven sampling stations are shown in Figure 1, along with the location of the release. The geodetic locations of these seven stations are summarized in Table 1. To simulate the activity concentration time series “seen” at these seven stations, we used a forward-time Lagrangian stochastic model driven by re-analyzed meteorological wind fields with a 6-h temporal resolution obtained over the temporal interval from 11 March 2011 00:00 Coordinated Universal Time (UTC) to 21 March 2011 12:00 UTC. The synthetic concentration time series at the seven stations were sampled with a temporal resolution of 12 h. The synthetic concentration data generated by the forward-time Lagrangian stochastic model (using the re-analyzed meteorological wind fields as input) were embedded within white and normally distributed noise with a standard deviation equal to 10% of the true concentration amplitude in order to simulate measurement uncertainty.

The adjunct concentration fields C^* required for the Bayesian computation were obtained using a backward-time LS model (which corresponds nominally to the adjoint of the forward-time LS model). Owing to errors in the numerical solution of the model equations that constitute the forward and backward LS models, it was found that duality relationship $\langle C|h \rangle = \langle C^*|S \rangle$ was not verified. The discrepancy in this relationship can be interpreted as the contribution to the model error component in the Bayesian analysis (viz., if the duality relationships were exactly satisfied in this example, then the uncertainty arising from the model error would vanish exactly). Indeed, informal tests conducted on a related LS dispersion model (utilizing

Table 1: The latitudinal and longitudinal coordinates of the seven sampling stations used for the emission rate profile reconstruction.

Location	Station ID	Latitude (°N)	Longitude (°E)
Sidney, BC	CAP14	48.6502	−123.399
Baja California	MXP44	29.9500	−115.117
Bilbino	RN057	68.0500	166.450
Petropavlovsk	RUP60	53.0167	158.650
Sacramento	USP70	38.5556	−121.469
Sand Point, Alaska	USP71	48.6502	−160.497
Midway Island	USP78	28.2000	−177.356

a Lagrangian particle-model kernel which is similar to that employed in the models applied herein) with respect to the forward-backward equivalence as embodied in the duality relationship has shown that the departure in this relationship is typically on the order of a factor of 2 or more.

Because the location of the source (Fukushima Daiichi nuclear power plant) was known *a priori* in this example, the source parameter vector here reduces to $\boldsymbol{\theta} = (Q_{s,1}, Q_{s,2}, \dots, Q_{s,N_s})$ with $N_s = 20$. Each component of $\boldsymbol{\theta}$ corresponds to the emission rate averaged over a 12-h time interval. More specifically, with reference to Eq. (20), the predicted concentration $\bar{C}_J(\boldsymbol{\theta})$ is obtained from

$$\bar{C}_J(\boldsymbol{\theta}) = \sum_{k=1}^{N_s} \bar{a}_{Jk} Q_{s,k}, \quad (27)$$

with

$$\bar{a}_{J,k} \equiv \int_{t_{k-1}}^{\min(t_J, t_k)} C^*(\mathbf{x}_s, t' | \mathbf{x}_J, t_J) dt', \quad (28)$$

where $\bar{a}_{Jk}$ is the coupling (or transfer) coefficient that relates the emission at the k -th time interval to the expected concentration measured at the space-time point $(\mathbf{x}_J, t_J)$. Note that in this case, the source location $\mathbf{x}_s$ is known *a priori* to be at the Fukushima Daiichi nuclear power plant (and, as a consequence, the values of the transfer coefficients $\bar{a}_{J,k}$ are known *a priori*).

In this example, we applied nested sampling to draw samples from the posterior distribution of $\boldsymbol{\theta}$ (source emission rate profile). Owing to the fact that the emission rate profile varies over many orders of magnitude, we reconstruct the common logarithm of the emission rate ($q \equiv \log_{10}(Q)$) rather than the emission rate directly. The prior distribution for q is assumed to be uniform (flat), implying a non-uniform (or non-flat) prior for Q . Indeed, recall that a uniform prior on a parameter set $\boldsymbol{\theta}$

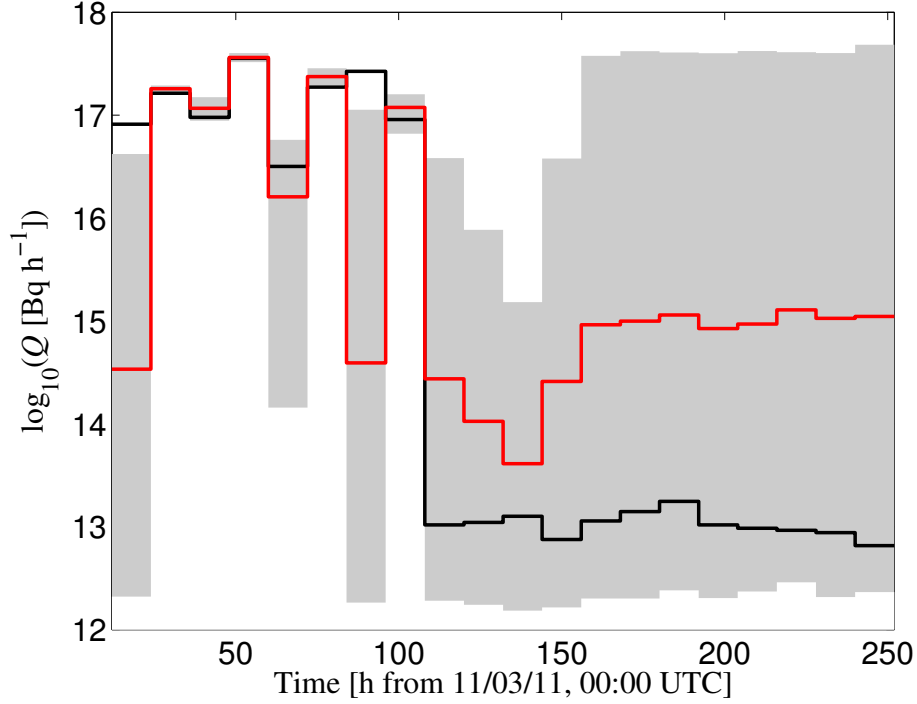


Figure 2: The best estimate for the Xe-133 emission rate time profile given by the posterior mean (solid red curve) compared to the true emission rate time profile (solid black curve), along with the 90% credible (HPD) interval for the emission rate profile estimate (gray-shaded region).

does not correspond to some uniform prior on some nonlinear function of it, $\mathcal{F}(\theta)$ (or vice-versa). The two priors are related by

$$p(\theta|I) = p(\mathcal{F}(\theta)|I) \left| \frac{\partial \mathcal{F}}{\partial \theta} \right|, \quad (29)$$

where $|\partial \mathcal{F} / \partial \theta|$ is the determinant for the Jacobian of the transformation. In the current example, logarithmic uniform priors are imposed on q with a lower bound of $q_{\min} \equiv \log_{10}(Q_{\min}) = 12$ and an upper bound of $q_{\max} = \log_{10}(Q_{\max}) = 18$ where Q has units of Bq h^{-1} . In the specification of the likelihood function [cf. Eq. (10)], the noise error variance is specified as $s_J^2 = s_{e,J}^2 + s_{m,J}^2$ where the measurement error standard deviation $s_{e,J}$ was assigned to be 10% of the measured concentration data d_J and the model error standard deviation $s_{m,J}$ was assigned (arbitrarily, or perhaps nominally) to be 25% of the predicted concentration $\bar{C}_J(\theta)$. Recall that the contribution to the model error here arises solely from numerical errors in the solution of the model equations for the forward and backward LS dispersion models which lead to the duality relationship not being verified exactly in practice (as it would be in theory in the absence of these numerical errors).

Figure 2 displays the best estimate of the time profile for the emission rate obtained

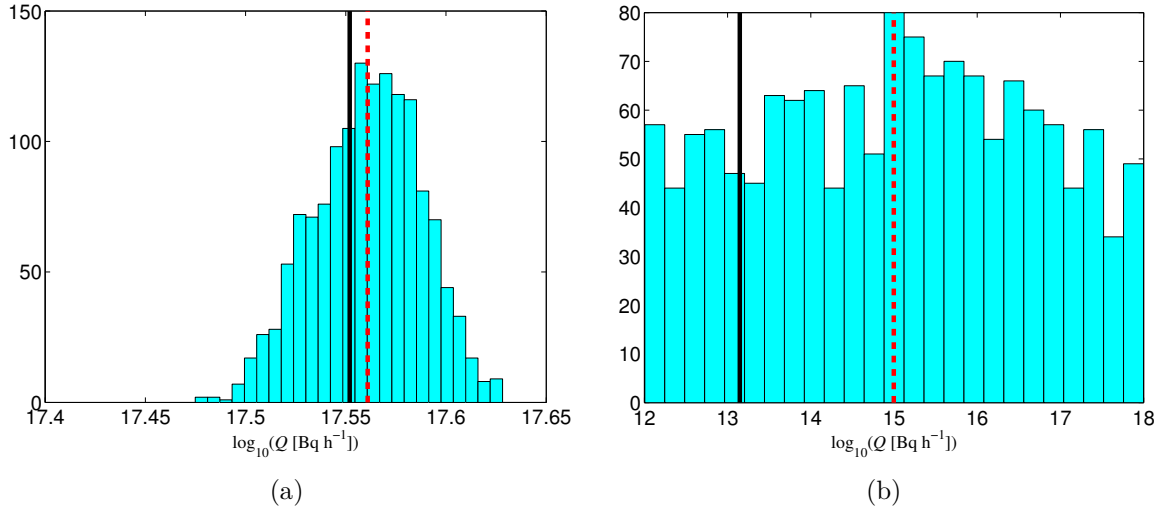


Figure 3: Marginal posterior probability distribution for $\log_{10}(Q)$ obtained for the 12-h time interval (a) (48,60) h and (b) (168,180) h after the arbitrary time origin t_0 (11 March 2011 00:00 UTC). The solid vertical line indicates the true value for the emission rate and the dashed vertical line corresponds to the best estimate of the emission rate obtained as the posterior mean of the associated marginal posterior probability distribution.

as the posterior mean (solid red curve) which can be compared with the true emission rate profile (solid black curve). In addition, the figure also presents the 90% credible (HPD) intervals for the emission rate profile (gray shaded area). It is noted that the gray-shaded region does appear to bracket the true emission rate profile at the indicated level of confidence. In general, the emission rates up to about 100 h (approximately or better) after the arbitrary time origin t_0 (taken to be 11 March 2011 00:00 UTC) are relatively well constrained by the observed activity concentration data (taken from the seven sampling stations). However, the 90% credible bounds for emission rates obtained for times greater than about 100 h after t_0 are seen to be very large, implying that the emission rates over this period of time are not well constrained by the information contained in the activity concentration time series.

The synthetic activity concentration time series data from the seven sampling stations *only* cover the time period up to 10 d after the arbitrary time origin t_0 . However, for many of the sampling sites in North America, it takes approximately 7 d after the initial release of radioxenon in a given (emission) period to be transported and dispersed across the Pacific Ocean by the prevailing winds and contribute to the measured concentration time series at a sampler location on the west coast of North America. As a consequence, the activity concentration time series from the majority of the sampling sites (which are truncated 10 d after the arbitrary time origin) do not contain contributions to the concentration time series arising from emissions at the Fukushima Daiichi plant that occur 100 h after t_0 . It is for this reason that the 90%

credible intervals for the emission rates are so large 100 h after t_0 , as these emission rates are not constrained by the concentration time series observed for the (restricted) time period of up to 10 d after t_0 .

Figure 3 shows the marginal posterior distribution of the emission rates for two 12-h time intervals after the time origin t_0 ; namely, for the interval (48, 60) h and (168, 180) h after t_0 . For the earlier time interval, it is seen that the marginal posterior distribution for the emission rate Q is well defined and occupies only a very small region interior to the logarithmic uniform prior distribution for this emission rate. This implies that the activity concentration data contains sufficient information to constrain this particular emission rate. Furthermore, note that the histogram for this case appears to be approximately Gaussian with the peak of the distribution centered on the true value of the emission rate for this interval. On the other hand, the marginal posterior distribution for the emission rate obtained at the later time interval is very broad, implying that the emission rate for this interval is not well constrained by the available activity concentration data (viz., the data does not contain sufficient information to estimate this parameter accurately). Indeed, the marginal posterior distribution for the logarithm of the emission rate for the later time interval is seen to coincide (approximately or better) with the log uniform prior distribution for this parameter, implying there is no information in the activity concentration that can be used to constrain this emission rate.

In addition, the limited temporal extent of the activity concentration time series used for the emission rate profile reconstruction in this example, demonstrates the advantage of the Bayesian inferential methodology. The credible intervals for the emission rate parameters shown in Figure 2 as well as the posterior distributions of the emission rate exhibited in Figure 3 provide an objective indication of certainty/uncertainty in the recovery of the emission rate. Without such knowledge, it is difficult to assess the significance that one should attribute to a particular estimate for the emission rate. This information allows one to assess objectively that the emission rates recovered at the later times are not constrained by the information in the activity concentration time series. To obtain better estimates for these later emission rates will require activity concentration time series be measured at times longer than 10 d after t_0 .

The nested sampling allows an estimation of the information gain $\mathcal{S}$ to be obtained. For this example, the information gain obtained from the activity concentration data $\mathbf{d}$ was found to be 18.6 natural units or nits (or, equivalently, 26.8 binary units or bits), implying that the information contained in the activity concentration allowed the “posterior volume” of the hypothesis space (volume of hypothesis space of reasonably large plausibility after the receipt of the concentration data) to decrease by a factor of $\exp(\mathcal{S}) \approx 1.194 \times 10^8$ relative to the “prior volume” of the hypothesis space (volume of hypothesis space of reasonably large plausibility before the receipt of the concentration data). Furthermore, this reduction in the hypothesis space does not

occur uniformly in every direction. The directions in the hypothesis space associated with the emission rates at the earlier times exhibit the most significant reduction in the “posterior volume” relative to the “prior volume”, whereas the directions in the hypothesis space associated with the emission rates at the later times do not exhibit a reduction of the posterior volume relative to the prior volume at all (cf. Figure 2).

5.2 Case 2: Source reconstruction

The current case study involves the use of some real measurements of activity concentrations obtained from the IMS radionuclide network for source reconstruction. The source for this example consists of the stack emissions from Chalk River Laboratories (CRL). Chalk River Laboratories houses an international production facility for medical radioisotopes. It is located at a latitude of 46.15° N and at a longitude of -77.37° E, about 180 km northwest of the city of Ottawa, Ontario. A characterization of the weekly stack emissions of Xe-133 from the CRL medical isotope production facility over a 5-year period yielded a median daily emission of about 24 TBq (or, equivalently, an emission rate of about 1.0×10^{12} Bq h⁻¹).

For this case study, we utilized Xe-133 activity concentration measurements obtained from three sampling sites in North America as shown in Figure 4. The three sites form part of the noble gas monitoring network of the International Monitoring System. In particular, the three stations exhibited in Figure 4 are as follows : CAX17 (St. John’s, Newfoundland located at latitude 47.59° N and longitude -52.74° E); USX75 (Charlottesville, Virginia located at latitude 38.0° N and longitude -78.4° E); and, USX74 (Ashland, Kansas located at latitude 31.17° N and longitude -99.77° E).

At each monitoring site, one of two different monitoring technologies is employed to measure radioxenon. The St. John’s site has a Système de Prélèvements et d’Analyse en Ligne d’Air pour quantifier le Xénon (SPALAX) high-resolution gamma system operating on a 24-h sample collection period, while the remaining sites have a Swedish Automatic Unit for Noble Gas Acquisition (SAUNA) beta-gamma coincidence system with a 12-h sample collection period. Both systems employ activated charcoal to remove xenon from the air for radioxenon analysis. After the measurement process is complete, the xenon sample can be stored in an archive bottle for optional re-measurement either on-site or in an off-site laboratory. The stable xenon volume collected varies approximately from between 2 ml to 8 ml depending on the technology used, but both technologies have a roughly equivalent Xe-133 sensitivity of about 0.2 mBq m⁻³.

Activity concentrations of Xe-133 for Case 2 were obtained for the single month of December 2011. For this case study, we used 36 activity concentration samples extracted from the three sampling sites. These concentration samples were then blocked-averaged over a temporal duration of 3 d to give 8 concentration data points



Figure 4: Locations of the three sampling stations (yellow markers) from the International Monitoring System radionuclide network used for Case 2. The location of the Xe-133 tracer source (red marker) was at Chalk River Laboratories.

d_J that were used as the input for the source reconstruction. In other words, $N = 8$ for this case study. As already mentioned, re-analyzed meteorological fields for December 2011 were used to “drive” an operational backward LS model to determine C^* . These adjunct concentration fields were then used for the rapid determination of the predicted concentration $\bar{C}_J(\theta)$ for an arbitrary source hypothesis θ .

For the source reconstruction, it is also necessary to provide an estimate for the noise error variance s_J^2 [cf. Eq. (10)]. This estimate includes two major contributions: (1) an estimate for the sensor error sampling variance $s_{e,J}^2$ in the measurement of d_J and (2) an estimate for the model error variance $s_{m,J}^2$ in the prediction of $\bar{C}_J(\theta)$. These two contributions to the noise error variance were added in quadrature to give $s_J^2 = s_{e,J}^2 + s_{m,J}^2$. Very good estimates for the sensor error standard deviation (or square root of the variance) were provided for the expected precision in the measurements of the activity concentration at each of the three sampling stations. In contrast, no estimates for the expected precision in the predicted concentrations were available. As a consequence, the model error standard deviation was assumed (rather arbitrarily) to be 50% of the predicted concentration $\bar{C}_J(\theta)$.

The emission source is near ground level so that the height of the source is not of any interest for the reconstruction. Consequently, the unknown source location parameters are taken to be the source longitudinal (x_s) and latitudinal (y_s) positions.

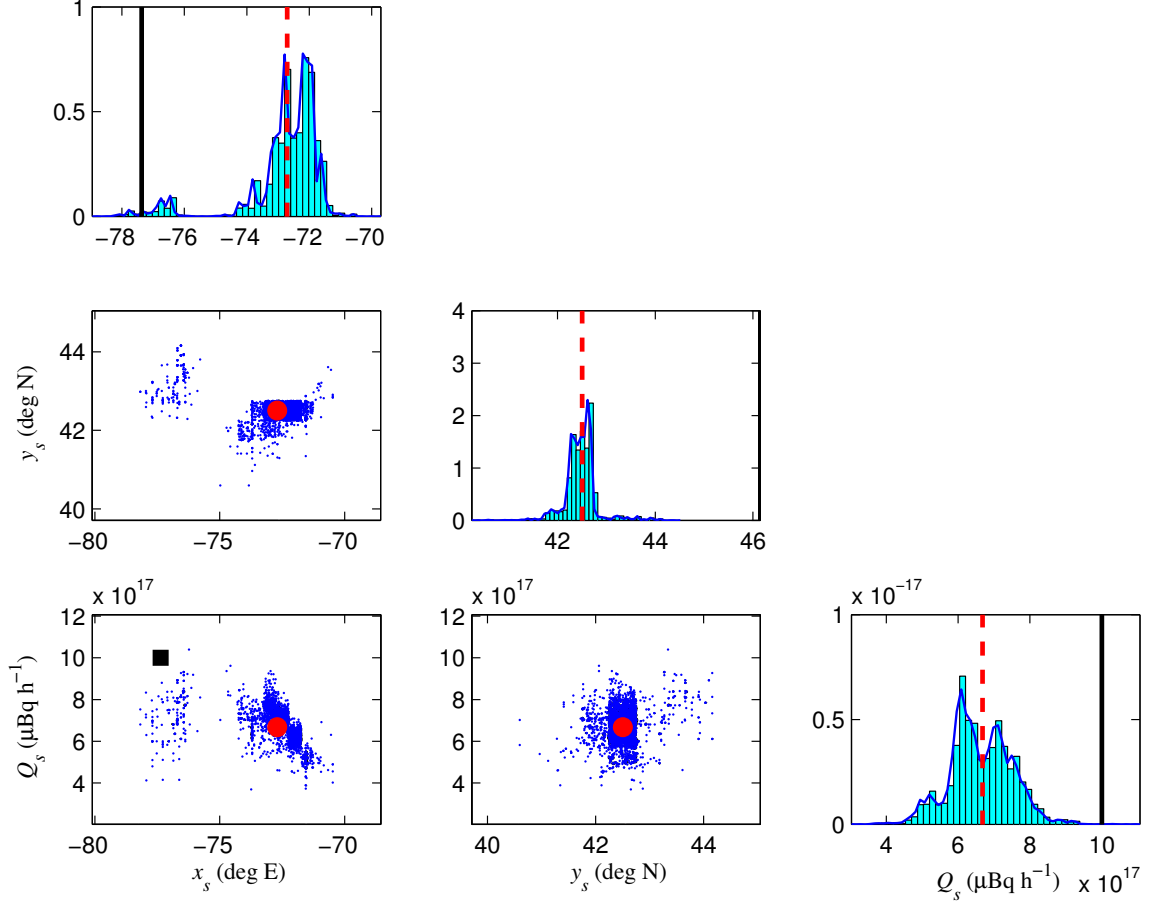


Figure 5: Univariate (diagonal) and bivariate (off-diagonal) marginal posterior distributions of the source parameters: longitudinal position x_s , latitudinal position y_s , and emission rate Q_s . A solid square or a solid vertical line indicates the true parameter values, whereas a solid circle or a dashed vertical line indicates a best estimate (posterior mean) for these parameter values.

Furthermore, the CRL release is approximated as a continuously emitting point source having a constant emission rate Q_s . In other words, $N_s = 1$ in Eq. (11) and $Q_s \equiv Q_{s,1}$ with $t_0 \rightarrow -\infty$ and $t_1 \rightarrow \infty$ for this case, so $\boldsymbol{\theta} = (x_s, y_s, Q_s)$. The MT-DREAM_(ZS) algorithm was used to draw samples from the posterior distribution of $\boldsymbol{\theta}$, with the initial population for the archive of past states obtained by sampling from the prior distribution $p(\boldsymbol{\theta}|I)$ [see Eq. (13)]. For this prior distribution $p(\boldsymbol{\theta}|I)$, the various hyperparameters are defined *a priori* as follows: (1) the minimum and maximum emission rates are prescribed as $Q_{\min} = 1.0 \times 10^{15} \mu\text{Bq h}^{-1}$ and $Q_{\max} = 1.0 \times 10^{20} \mu\text{Bq h}^{-1}$, respectively and (2) the prior bounds for the source location (x_s, y_s) are given by $\mathcal{D} = (-125^\circ \text{ E}, -45^\circ \text{ E}) \times (25^\circ \text{ N}, 75^\circ \text{ N})$ (an *a priori* spatial domain that encompasses the entire North American continent).

Table 2: The posterior mean, posterior standard deviation, and lower and upper bounds of the 95% HPD interval of the parameters x_s ($^{\circ}\text{E}$), y_s ($^{\circ}\text{N}$), and Q_s ($\mu\text{Bq h}^{-1}$) calculated from samples of θ drawn from the posterior distribution $p(\theta|\mathbf{d}, \mathbf{s}, \alpha, \beta, I)$.

Parameter	Mean	Standard Deviation	95% HPD	Actual
x_s ($^{\circ}\text{E}$)	-72.71	1.22	(-76.71, -71.56)	-77.37
y_s ($^{\circ}\text{N}$)	42.51	0.32	(41.86, 43.35)	46.15
Q_s ($\mu\text{Bq h}^{-1}$)	6.68×10^{17}	8.55×10^{16}	$(4.98, 8.33) \times 10^{17}$	1.0×10^{18}

The samples of θ drawn from the posterior distribution were used to determine the source characteristics (location, emission rate). To this purpose, source parameter estimates were obtained directly from the multiple chains of samples generated by the MT-DREAM_(ZS) algorithm after convergence of the chains as determined by the Gelman-Rubin $\hat{R}$ statistic. Figure 5 displays the univariate (diagonal) and bivariate (off-diagonal) marginal posterior distributions for the source parameters. The solid vertical line delineates the true value of the parameter and the dashed vertical line indicates the best estimate (posterior mean) of the parameter in the univariate marginal posterior distributions. Similarly, for the bivariate marginal posterior distribution of various combinations of parameters, the solid square marks the position of the true parameter values and the solid circle exhibits the best estimates (posterior means) of these parameters. It should be noted that the axes limits are chosen to display the regions of highest probability in the marginal posterior probability distributions for the parameters. As a result, in certain cases these regions do not contain the true parameter values (with the result that some of the panels in Figure 5 do not contain either the solid square or solid vertical line representing the true parameter values).

Table 2 summarizes the posterior mean, posterior standard deviation, and lower and upper bounds for the 95% credible (or HPD) interval of the recovered source parameters. A perusal of these values shows that the *accuracy* in the source reconstruction is quite good for both the location (considering the fact that the reconstruction was undertaken on a continental scale) and the emission rate. In particular, the distance between the true source location and its estimate (obtained as the posterior mean) is only about 572 km (see Figure 6). Finally, the best estimate of the emission rate (obtained again as the posterior mean) is within 33% of the true emission rate.

However, an examination of Table 2 shows that the *precision* in the source parameter estimates is poorly determined. More specifically, the 95% HPD intervals for the longitudinal and latitudinal source positions and for the emission rate do not contain the true source parameters. This defect in the source reconstruction can be attributed to the difficulties in providing good estimates for the model errors in the prediction of $\bar{C}_J(\theta)$. This inability to provide good estimates for the model errors leads to

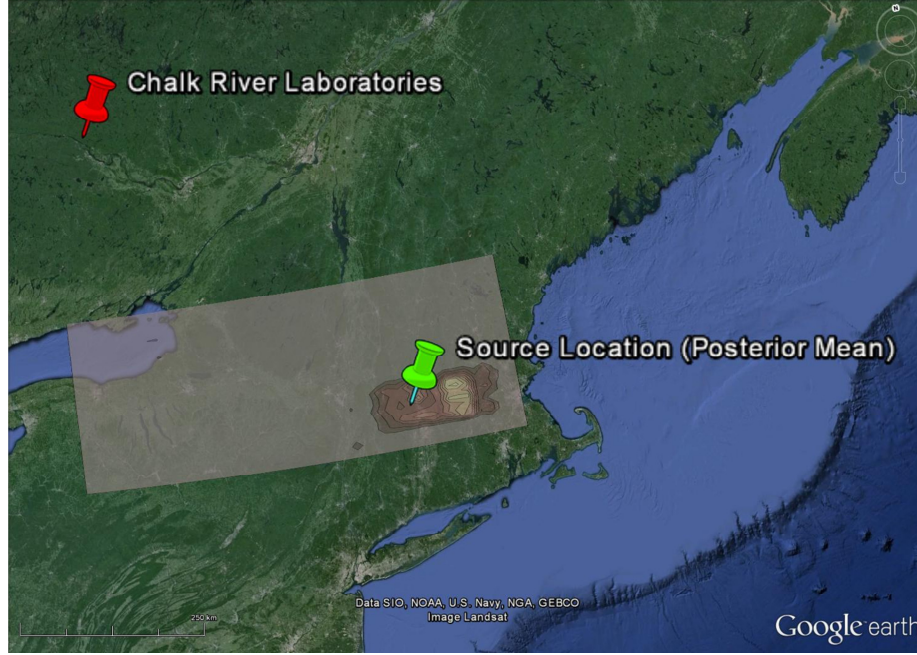


Figure 6: Two-dimensional marginal posterior distribution of the source location geo-referenced on a Google Earth image.

a loss of power in the source reconstruction, and may mask important features of the measured activity concentration data. Finally, it should be noted that unlike Case 1 described above where the model errors in the synthetic (albeit idealized) concentration data arose only from numerical errors in the solution of the model equations for the forward-time and backward-time LS models, the model errors in Case 2 for the real concentration data may arise from numerous sources (as outlined briefly in Section 3.1) and can be much more significant.

Given the complexity in the heteroscedastic variance of the model error, its *a priori* specification is highly problematic. In light of these difficulties, let us consider an alternative measurement model to that introduced earlier in Eq. (6). To this end, let us now focus on the following alternative measurement model:

$$d_J = m_J \bar{C}_J(\boldsymbol{\theta}) + n_J, \quad J = 1, 2, \dots, N, \quad (30)$$

where m_J are unknown multipliers (scale factors) that are applied to the predicted concentration $\bar{C}_J(\boldsymbol{\theta})$ in order to compensate for the model uncertainty. The n_J term in Eq. (30) represents *only* the measurement error in the activity concentration d_J . The model error arising from the predicted concentration $\bar{C}_J(\boldsymbol{\theta})$ is compensated through the use of the multipliers m_J ($J = 1, 2, \dots, N$).

For the alternative measurement model, the multipliers m_J ($J = 1, 2, \dots, N$) are unknown parameters that need to be estimated in addition to the usual source param-

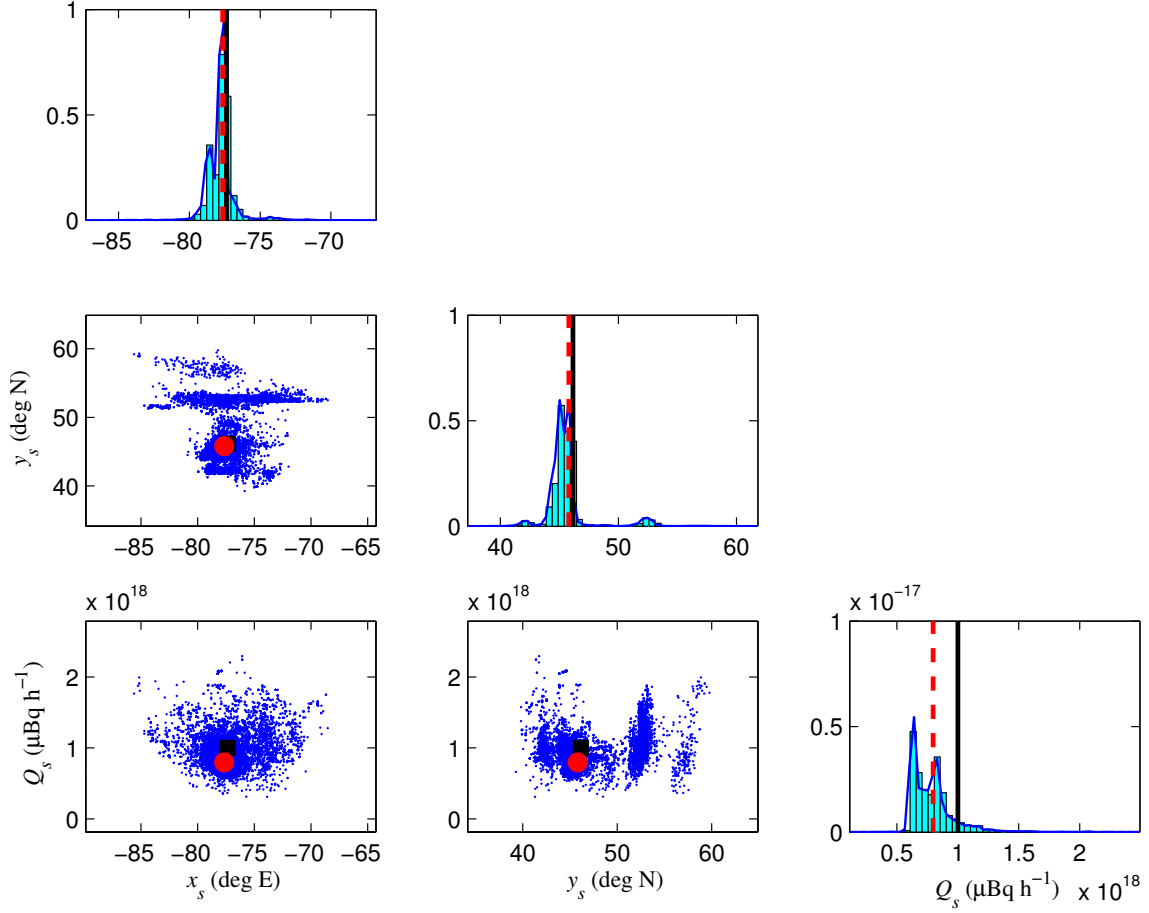


Table 3: The posterior mean, posterior standard deviation, and lower and upper bounds of the 95% HPD interval of the parameters x_s ($^{\circ}\text{E}$), y_s ($^{\circ}\text{N}$), and Q_s ($\mu\text{Bq h}^{-1}$) calculated from samples of $\boldsymbol{\theta}$ drawn from the posterior distribution $p(\boldsymbol{\theta}|\mathbf{d}, \mathbf{s}, \alpha, \beta, I)$. The results are obtained using the alternative measurement model.

Parameter	Mean	Standard Deviation	95% HPD	Actual
x_s ($^{\circ}\text{E}$)	-77.66	1.04	(-79.21, -74.88)	-77.37
y_s ($^{\circ}\text{N}$)	45.80	1.96	(42.69, 52.64)	46.15
Q_s ($\mu\text{Bq h}^{-1}$)	7.97×10^{17}	1.75×10^{17}	$(6.37, 12.5) \times 10^{17}$	1.0×10^{18}

ters. Let us denote the source parameters in this alternative model by $\boldsymbol{\theta}^s \equiv (\mathbf{x}_s, Q_s)$. Furthermore, let $\boldsymbol{\theta}^m \equiv (m_1, m_2, \dots, m_N)$ denote all other relevant parameters (multipliers in our example). These latter parameters are referred to as nuisance parameters. Both sets of parameters define the parameter vector as $\boldsymbol{\theta} = (\boldsymbol{\theta}^s, \boldsymbol{\theta}^m)$. The likelihood function for the alternative measurement model is still given by Eq. (10). However, for this case the estimated noise variances s_J^2 appearing in the likelihood function now *only* include the measurement error contribution. In other words, $s_J^2 = s_{e,J}^2$. As already noted above, this contribution to the uncertainty is well characterized in our current application implying that the prior uncertainty in the measurement errors can be specified correctly. For the alternative measurement model, we need to specify also the prior distributions for the multipliers (nuisance parameters). Towards this end, uniform priors defined over the range $(m_{\min}, m_{\max})$ will be used as priors for the multipliers. Consequently, the prior distribution for the alternative measurement model replaces Eq. (13) by the following assignment (recalling again that $N_s = 1$ and $Q_s \equiv Q_{s,1}$ for Case 2):

$$p(\boldsymbol{\theta}|I) \propto \mathbb{I}_{\mathcal{D}}(\mathbf{x}_s) \mathbb{I}_{(Q_{\min}, Q_{\max})}(Q_s) \times \prod_{J=1}^N \mathbb{I}_{(m_{\min}, m_{\max})}(m_J). \quad (31)$$

Using the alternative measurement model and the modified likelihood function and prior distribution, we applied the MT-DREAM_(ZS) algorithm to sample from the modified posterior distribution for $\boldsymbol{\theta}$. The hyperparameters that define the prior distribution for the source parameters $\boldsymbol{\theta}^s$ were exactly as described above. The lower and upper prior bounds for the multipliers m_J were $m_{\min} = 0.1$ and $m_{\max} = 10.0$, respectively ($J = 1, 2, \dots, N$ where $N = 8$). The univariate and bivariate marginal posterior distributions of the various source parameters are displayed in Figure 7. The solid square or solid vertical line marks the true parameter values. These should be compared with their best estimates (posterior means) marked by either a solid circle or a dashed vertical line. Table 3 summarizes the posterior mean, posterior standard deviation, and lower and upper bounds for the 95% HPD interval for the source parameters.

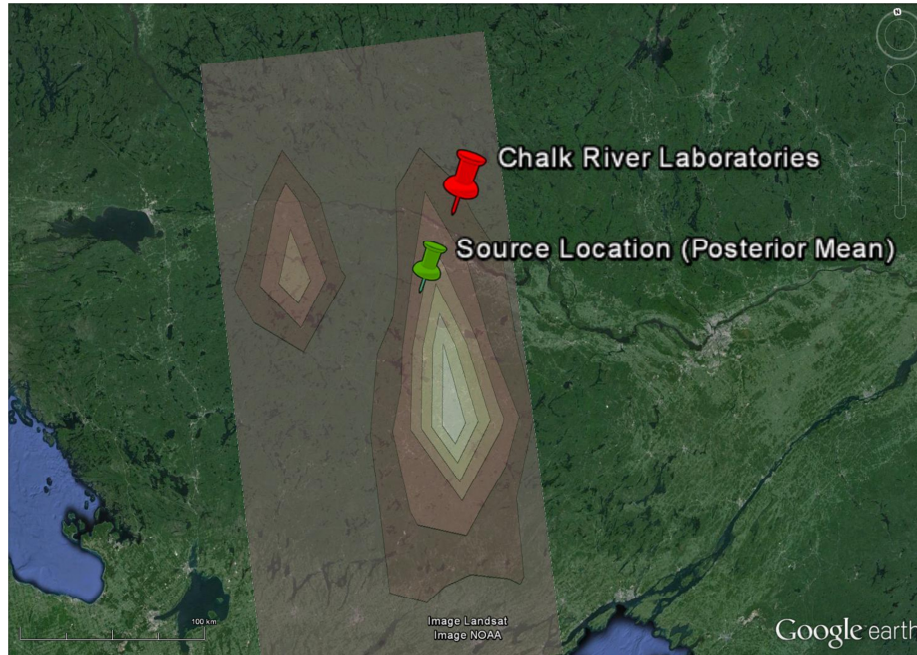


Figure 8: Two-dimensional marginal posterior distribution of the source location geo-referenced on a Google Earth image. The marginal distribution is for the alternative measurement model.

A perusal of Figure 7 and of Table 3 shows that all the source parameters have been recovered with very good accuracy. Indeed, the distance between the actual source location and the best estimate (posterior mean) of the source location is about 44 km (see Figure 8), and the recovered emission rate is within 20% of the actual emission rate. The accuracy in the inferred source location is roughly a factor of ten better than that obtained using the standard measurement model (which does not use multipliers to try to compensate for the unknown model errors). More importantly, the precision of the source parameter estimates (viz., the 95% credible intervals) for the alternative measurement model contain the actual (true) values for the source parameters, in stark contrast to the reconstruction using the standard measurement model. This is evident if we compare the inferred source parameter values and their uncertainty bounds in Table 2 with those in Table 3.

In the practical real-world it is difficult to correctly specify *a priori* the structure and scale of the model error. In view of this, the incorporation of multipliers with the predicted concentrations for model error compensation can improve significantly the quality of the source reconstruction. Figure 9 shows the univariate marginal posterior distributions for the various multipliers m_J ($J = 1, 2, \dots, 8$) and provides evidence of the complexity in the model error structure for the current problem. Indeed, a perusal of this figure readily makes evident the complexity in the heteroscedastic model error structure. Note that the distributions for the multipliers associated with some of the

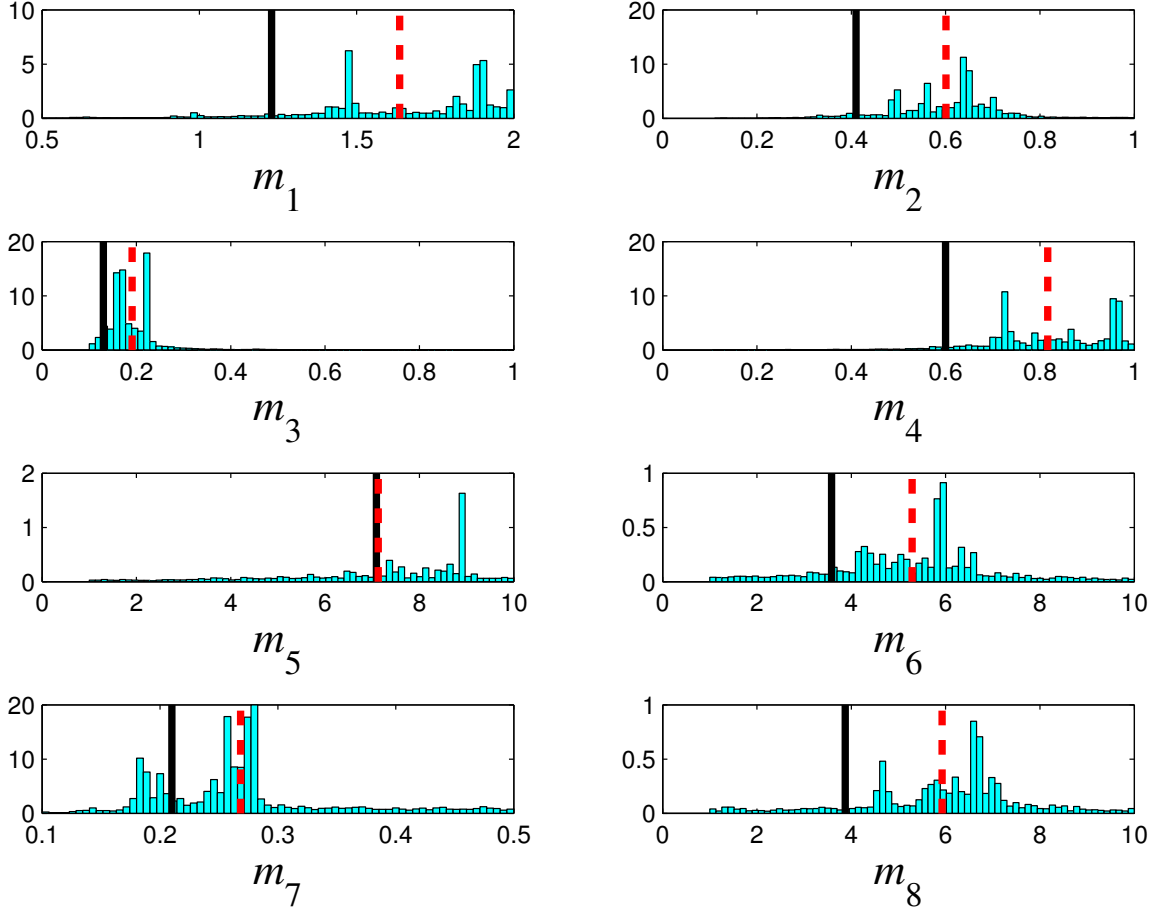


Figure 9: Univariate marginal posterior distribution of the multipliers m_J ($J = 1, 2, \dots, 8$). The “actual” values of the multipliers are indicated with the solid vertical line which should be compared with the best estimates (posterior means) of the multipliers.

predicted concentrations are quite broad, indicating that the model error for these predicted concentrations are not well determined. In other cases, the distributions for the multipliers (e.g., m_2 and m_3) are quite narrow, implying that the data contain reasonable information to constrain the model errors associated with these multipliers fairly well.

Although the true values of the multipliers are unknown, we can nevertheless obtain an independent estimate for their values as follows. We can use the actual source parameters θ_s^* to predict the activity concentration $\bar{C}_J(\theta_s^*)$ that would be expected at the various observation stations. Using this information, we can provide an independent point estimate for the multipliers as $\hat{m}_J = d_J / \bar{C}_J(\theta_s^*)$, $J = 1, 2, \dots, 8$. The solid vertical lines in Figure 9 exhibit these “actual” values for the multipliers. These values should be compared with best estimates (posterior means) of the multipliers. Observe that the best estimates of the multipliers are broadly consistent with the

“actual” values for the multipliers obtained from the indicated point estimates. Furthermore, the width of the posterior distribution of the various multipliers are seen to enclose the “actual” values for the multipliers. Finally, even though some of the multipliers are highly uncertain, inclusion of multipliers to compensate for the model errors does lead to significantly improved estimates for the source parameters (both in accuracy and in precision).

6 Discussion and conclusions

In this study, a Bayesian inferential methodology (Bayesian probability theory) has been proposed for source reconstruction in the context of activity concentration data that would be available from an worldwide operational network of sensors (International Monitoring System). This methodology for source reconstruction has been applied to two different case studies: namely, emission rate profile reconstruction for the Fukushima Daiichi nuclear power plant release using synthetic activity concentration data generated for seven sampling stations and a more difficult situation involving the reconstruction of the source characteristics from the CRL release using only a small number of real (actual) activity concentration measurements (8 measurements) obtained from only three sampling stations. All these stations were part of the IMS radionuclide network. Both case studies involved the utilization of an operational backward LS model for long-range transport by the atmosphere on a continental or hemispheric scale.

A detailed description of Bayesian probability theory as applied specifically to source reconstruction has been provided in this report. Furthermore, efficient and robust MCMC algorithms (e.g., nested sampling, multiple-try differential evolution adaptive Metropolis sampling) for estimating and summarizing the information embodied in the posterior probability distribution of the source parameters θ have been described and applied successfully to the two case studies chosen to illustrate the utility of the proposed methodology. More specifically, it has been demonstrated how Bayesian probability theory in conjunction with MCMC provides both a rigorous theoretical and practically applicable framework for the recovery of various source parameters (e.g., location, emission rates) and for quantification of the uncertainties in the recovered estimates that take into account the measurement errors in the concentration data, the model errors in the predicted concentrations used for the interpretation of the data, and any prior information we may have.

The experience with addressing the second case involving the use of real activity concentration data demonstrated that the principal difficulty in the reconstruction lay in providing the correct *a priori* specification of the model error for the various predicted concentrations associated with the measured concentrations used for the source parameter recovery. A naïve specification for the model error gave a source

reconstruction that was quite reasonable in terms of the accuracy in the recovery of the location and emission rate, but the precision in the estimates was generally poor in the sense that the reported uncertainty bounds in the recovery of the source parameters did not include the actual (true) values for these parameters. This led to the development of an alternative measurement model which incorporated multipliers (scale factors) with the predicted concentrations in order to compensate for the model errors. The application of this alternative model was shown to yield significantly improved estimates for the source parameters, both in terms of their accuracy and their precision. Undoubtedly more sophisticated measurement models (required for the assignment of the likelihood function) can be formulated, being purely a matter of improved intuition and educated intelligence (creativity, insight and imagination), which can lead to further improvements in the source reconstruction.

The avalanche of new data obtained from the International Monitoring System that is managed under the auspices of the United Nations Comprehensive Nuclear-Test-Ban Treaty Organization for treaty compliance verification will require ever more sophisticated statistical tools for source reconstruction. In this report, we describe a preliminary application of Bayesian probability theory for source reconstruction in the context of using IMS station data for possible compliance verification. Even so, the current effort has shown that its successful application to real-world problems using real sensor networks and operational dispersion models will require a better understanding of both the scale and structure of the model error in the predicted concentrations. It is anticipated that the proper incorporation of information about the underlying model error in conjunction with Bayesian inference employing state-of-the-science MCMC sampling algorithms (such as nested sampling and differential evolution adaptive Metropolis sampling) can potentially provide a very flexible framework for source reconstruction.

The next step in this effort would be to provide more validation of the proposed Bayesian inference methodology for source reconstruction. This would involve the careful consideration of additional cases that embody more complex source distributions and situations that are of relevance to the detection and identification of nuclear detonation events by the observational technology available in the IMS. Furthermore, it would be useful to generalize the Bayesian methodology described herein in the following manner. The focus in this report has been on the application of Bayesian probability theory to the parametric estimation of various source parameters θ using only a single atmospheric dispersion model for the predicted concentration. However, situations can arise when one has available several models for the predicted concentration, each of which can be driven by one or more (different) re-analyzed or forecasted meteorological wind fields. The question that arises is how to combine these various concentration predictions from the various models to improve the inference of the source parameters θ . It is anticipated that this important problem can be addressed in a straightforward manner as the formal Bayesian approach already

encompasses the rigorous schema for utilizing multimodel ensembles for inference. Indeed, a straightforward generalization of the inference methodology to accommodate multimodel ensembles through a model-averaged posterior distribution would enable one to combine uncertainty at both the model and parameter levels in a rigorous fashion and also to fuse different diagnostic data sets (e.g., radionuclide, seismic, infrasound)¹ in the parameter space, providing ultimately a more complete description of the state of knowledge of θ .

¹It should be noted that the International Monitoring System also includes seismological, hydroacoustical, and infrasound diagnostics which can be potentially used in addition to the diagnostic (activity concentration) provided by the radiological portion of the IMS for source term estimation.

This page intentionally left blank.

References

- [1] Wotawa, G., Becker, A., Kalinowski, M., Saey, P., Tuma, M., and Zähringer, M. (2010), Computation and analysis of the global distribution of radioxenon isotope ^{133}Xe based on emissions from nuclear power plants and radioisotope production facilities and its relevance for the verification of the Nuclear-Test-Ban Treaty, *Atmospheric and Chemical Physics*, 167, 541–557.
- [2] Hoffman, I., Ungar, K., Bean, M., Yi, J., Servranckx, R., Zaganescu, C., Ek, N., Blanchard, X., Le Petit, G., Brachet, G., Achim, P., and Taffary, T. (2009), Changes in radioxenon observations in Canada and Europe during medical isotope production facility shutdown in 2008, *Journal of Radioanalytical and Nuclear Chemistry*, 282, 767–772.
- [3] Wotawa, G., De Geer, L. E., Denier, P., Kalinowski, M., Toivonen, H., D’Amours, R., Desiato, F., Issartel, J. P., Langer, M., Seibert, P., Frank, A., Sloan, C., and Yamazawa, H. (2003), Atmospheric transport modeling in support of CTBT verification – overview and basic concepts, *Atmospheric Environment*, 37, 2529–2537.
- [4] Ringbom, A., Axelsson, A., Aldener, M., Auer, M., Bowyer, T., Fritioff, T., Hoffman, I., Khrustalev, K., Nikkinen, M., Popov, V., Popov, Y., Ungar, K., and Wotawa, G. (2014), Radioxenon detections in the CTBT International Monitoring System likely related to the announced nuclear test in North Korea on February 12, 2013, *Journal of Environmental Radioactivity*, 128, 47–63.
- [5] Yee, E. (2005), *Probabilistic Inference: An Application to the Inverse Problem of Source Function Estimation*, The Technical Cooperation Program (TTCP) Chemical and Biological Defence (CBD) Group Technical Panel 9 (TP-9) Annual Meeting, Defence Science and Technology Organization, Melbourne, Australia.
- [6] Yee, E. (2006), *A Bayesian Approach for Reconstruction of the Characteristics of a Localized Pollutant Source from a Small Number of Concentration Measurements Obtained by Spatially Distributed “Electronic Noses”*, Russian-Canadian Workshop on Modeling of Atmospheric Dispersion Weapon Agents, Karpov Institute of Physical Chemistry, Moscow, Russia.
- [7] Keats, A., Yee, E., and Lien, F.-S. (2007), Bayesian inference for source determination with application to a complex environment, *Atmospheric Environment*, 41, 465–479.
- [8] Chow, F., Kosovic, B., and Chan, S. (2008), Source inversion for contaminant plume dispersion in urban environments using building-resolving simulations, *Journal of Applied Meteorology and Climatology*, 47, 1553–1572.

- [9] Keats, A., Yee, E., and Lien, F.-S. (2007), Efficiently characterizing the origin and decay rate of a nonconservative scalar using probability theory, *Ecological Modeling*, 205, 437–452.
- [10] Keats, A., Yee, E., and Lien, F.-S. (2010), Information-driven receptor placement for contaminant source determination, *Environmental Modeling Software*, 25, 1000–1013.
- [11] Yee, E., Lien, F.-S., Keats, A., and D’Amours, R. (2008), Bayesian inversion of concentration data: Source reconstruction in the adjoint representation of atmospheric diffusion, *Journal of Wind Engineering and Industrial Aerodynamics*, 96, 1805–1816.
- [12] Yee, E. (2008), Theory for reconstruction of an unknown number of contaminant sources using probabilistic inference, *Boundary-Layer Meteorology*, 127, 359–394.
- [13] Yee, E. (2012), Probability theory as logic: Data assimilation for multiple source reconstruction, *Pure and Applied Geophysics*, 169, 499–517.
- [14] Yee, E. (2012), Source reconstruction: A statistical mechanics perspective, *International Journal of Environment and Pollution*, 48, 203–213.
- [15] Yee, E. and Flesch, T. K. (2010), Inference of emission rates from multiple sources using Bayesian probability theory, *Journal of Environmental Monitoring*, 12, 622–634.
- [16] Yee, E. (2012), Inverse dispersion for an unknown number of sources: Model selection and uncertainty analysis, *ISRN Applied Mathematics*, 2012, Article ID 465320, 20 pp.
- [17] CTBTO Preparatory Commission (2001), *The Global Verification Regime and the International Monitoring System*, Preparatory Commission for the Comprehensive Nuclear-Test-Ban Treaty Organization (CTBTO), Vienna.
- [18] Yee, E., Hoffman, I., and Ungar, K. (2014), Bayesian inference for source reconstruction: A real-world example, *ISRN Research Notices*, 2014, Article ID 507634, 12 pp.
- [19] Cox, R. T. (1946), Probability, frequency, and reasonable expectation, *American Journal of Physics*, 14, 1–13.
- [20] Jaynes, E. T. (2003), *Probability Theory: The Logic of Science*, Cambridge University Press, Cambridge, UK.
- [21] Cover, T. M. and Thomas, J. A. (1991), *Elements of Information Theory*, John-Wiley and Sons, Inc., New York.

- [22] Higdon, D., Gattiker, J., Williams, B., and Rightley, M. (2008), Computer model calibration using high-dimensional output, *Journal of the American Statistical Association*, 103, 570–583.
- [23] Gilks, W. R., Richardson, S., and Spiegelhalter, D. J. (1996), *Markov chain Monte Carlo in Practice*, CRC Press, Chapman and Hall.
- [24] Gelman, A., Carlin, J., Stern, H., and Rubin, D. (2003), *Bayesian Data Analysis*, Second Edition (Texts in Statistical Science), CRC Press, Chapman and Hall.
- [25] Laloy, E. and Vrugt, J. A. (2012), High-dimensional posterior exploration of hydrologic models using multiple-try DREAM_(ZS) and high-performance computing, *Water Resources Research*, 48, W01526, 18 pp.
- [26] Liu, J. S., Liang, F., and Wong, W. H. (2000), The multiple-try method and local optimization in Metropolis sampling, *Journal of the American Statistical Association*, 95, 121–134.
- [27] Gelman, A. G. and Rubin, D. B. (1992), Inference from iterative simulation using multiple sequences, *Statistical Science*, 7, 457–472.
- [28] Skilling, J. (2006), Nested sampling for general Bayesian computation, *Bayesian Analysis*, 1, 833–860.
- [29] Castillo, E. (1988), *Extreme Value Theory in Engineering*, Academic Press, Inc., San Diego.
- [30] Feroz, F., Hobson, M. P., and Bridges, M. (2009), MultiNest: an efficient and robust Bayesian inference tool for cosmology and particle physics, *Monthly Notices of the Royal Astronomical Society*, 398, 1601–1614.
- [31] Amano, H., Akiyama, M., Chunlei, B., Kawamura, T., Kishimoto, T., Kuroda, T., Muroi, T., Odaira, T., Ohta, Y., Takeda, K., Watanabe, Y., and Morimoto, T. (2012), Radiation measurements in the Chiba metropolitan area and radiological aspects of fallout from the Fukushima Daiichi nuclear power plant accidents, *Journal of Environmental Radioactivity*, 111, 42–52.
- [32] Stohl, A., Seibert, P., Wotawa, G., Arnold, D., Burkhart, J. F., Eckhardt, S., Tapia, C., Vargas, A., and Yasunari, T. J. (2012), Xenon-133 and caesium-137 releases into the atmosphere from the Fukushima Daiichi nuclear power plant: determination of the source term, atmospheric dispersion, and deposition, *Atmospheric and Chemical Physics*, 12, 2313–2343.

This page intentionally left blank.

DOCUMENT CONTROL DATA		
(Security markings for the title, abstract and indexing annotation must be entered when the document is Classified or Designated.)		
1. ORIGINATOR (The name and address of the organization preparing the document. Organizations for whom the document was prepared; e.g., Centre sponsoring a contractor's report, or tasking agency, are entered in section 8.) DRDC – Suffield Research Centre Box 4000, Station Main, Medicine Hat AB T1A 8K6, Canada		2a. SECURITY MARKING (Overall security marking of the document, including supplemental markings if applicable.) UNCLASSIFIED
		2b. CONTROLLED GOODS (NON-CONTROLLED GOODS) DMC A REVIEW: GCEC APRIL 2011
3. TITLE (The complete document title as indicated on the title page. Its classification should be indicated by the appropriate abbreviation (S, C or U) in parentheses after the title.) Bayesian inference for source term estimation: Application to the International Monitoring System radionuclide network		
4. AUTHORS (Last name, followed by initials – ranks, titles, etc. not to be used.) Yee, E.; Hoffman, I.; Ungar, K.; Malo, A.; Ek, N.; Bourgouin, P.		
5. DATE OF PUBLICATION (Month and year of publication of document.) October 2014	6a. NO. OF PAGES (Total containing information. Include Annexes, Appendices, etc.) 54	6b. NO. OF REFS (Total cited in document.) 32
7. DESCRIPTIVE NOTES (The category of the document, e.g., technical report, technical note or memorandum. If appropriate, enter the type of report, e.g., interim, progress, summary, annual or final. Give the inclusive dates when a specific reporting period is covered.) Scientific Report		
8. SPONSORING ACTIVITY (The name of the department project office or laboratory sponsoring the research and development – include address.) DRDC – Suffield Research Centre Box 4000, Station Main, Medicine Hat AB T1A 8K6, Canada		
9a. PROJECT OR GRANT NO. (If appropriate, the applicable research and development project or grant number under which the document was written. Please specify whether project or grant.) CSSP-2013-CD-1132	9b. CONTRACT NO. (If appropriate, the applicable number under which the document was written.)	
10a. ORIGINATOR'S DOCUMENT NUMBER (The official document number by which the document is identified by the originating activity. This number must be unique to this document.) DRDC-RDDC-2014-R71	10b. OTHER DOCUMENT NO(s). (Any other numbers which may be assigned this document either by the originator or by the sponsor.)	
11. DOCUMENT AVAILABILITY (Any limitations on further dissemination of the document, other than those imposed by security classification.) ☒ Unlimited distribution ☐ Defence departments and defence contractors; further distribution only as approved ☐ Defence departments and Canadian defence contractors; further distribution only as approved ☐ Government departments and agencies; further distribution only as approved ☐ Defence departments; further distribution only as approved ☐ Other (please specify):		
12. DOCUMENT ANNOUNCEMENT (Any limitation to the bibliographic announcement of this document. This will normally correspond to the Document Availability (11). However, where further distribution (beyond the audience specified in (11)) is possible, a wider announcement audience may be selected.)		

13. ABSTRACT (A brief and factual summary of the document. It may also appear elsewhere in the body of the document itself. It is highly desirable that the abstract of classified documents be unclassified. Each paragraph of the abstract shall begin with an indication of the security classification of the information in the paragraph (unless the document itself is unclassified) represented as (S), (C), or (U). It is not necessary to include here abstracts in both official languages unless the text is bilingual.)

In recent years, there has been an enormous quantity of data obtained from the International Monitoring System radionuclide network for the verification of the Comprehensive Nuclear-Test-Ban Treaty. The complexity of the instruments deployed here, of the radionuclide sources, and of the myriad of scientific questions related to treaty verification lead invariably to complex inference problems (associated with source term estimation) that require the application of sophisticated statistical tools. In this report, we demonstrate that a rigorous and general framework for addressing these problems is through Bayesian probability theory, allowing the rational inference of the posterior probability distribution of the source parameters of interest given any prior information and available activity concentration measurements. The methodology is demonstrated by application to two different problems: namely, the emission rate profile reconstruction of a radioxenon release from the Fukushima Daiichi nuclear power plant and source reconstruction (location and emission rate) of a radioxenon release from the Chalk River Laboratories (CRL) medical isotope production facility. The sampling of the resulting posterior distribution of the source parameters is undertaken using two different Markov chain Monte Carlo techniques: namely, nested sampling and multiple-try differential evolution adaptive Metropolis sampling with a past archive.

For the Fukushima nuclear power plant release, it is demonstrated that the *limited* temporal extent of the activity concentration time series obtained from the seven sampling sites cannot be used to constrain the emission rate profile at a later time. In particular, the Bayesian credible intervals for the reconstruction of the emission rate profile provide a quantitative indication of the uncertainty in this quantity, allowing an objective assessment of the fact that the emission rates recovered at the later times are not constrained by the information in the available activity concentration data. For the CRL release, it is shown that the principal difficulty in the reconstruction lay in the correct specification of both the scale and structure of the model errors used in the Bayesian inferential methodology. Consequently, for this case, two different measurement models for incorporation of the model errors in the predicted concentrations are considered. The performance of both of these measurement models with respect to their accuracy and precision in the recovery of the source parameters is compared and contrasted. In particular, it is shown that the incorporation of multipliers in the measurement model to compensate for the unknown model errors lead to significantly improved estimates for the source parameter (both in accuracy and in precision) compared to the simpler measurement model which does not use multipliers.

14. KEYWORDS, DESCRIPTORS or IDENTIFIERS (Technically meaningful terms or short phrases that characterize a document and could be helpful in cataloguing the document. They should be selected so that no security classification is required. Identifiers, such as equipment model designation, trade name, military project code name, geographic location may also be included. If possible keywords should be selected from a published thesaurus. e.g. Thesaurus of Engineering and Scientific Terms (TEST) and that thesaurus identified. If it is not possible to select indexing terms which are Unclassified, the classification of each should be indicated as with the title.)

atmospheric dispersion; Bayesian inference; International Monitoring System radionuclide network; Markov chain Monte Carlo; model and measurement errors; source reconstruction

DRDC | RDDC

SCIENCE, TECHNOLOGY AND KNOWLEDGE
FOR CANADA'S DEFENCE AND SECURITY

SCIENCE, TECHNOLOGIE ET SAVOIR
POUR LA DÉFENSE ET LA SÉCURITÉ DU CANADA



www.drdc-rddc.gc.ca