

A Simulation-Based Geostatistical Approach to Real-Time Reconciliation of the Grade Control Model

Wambeke, T.; Benndorf, J.

DOI

[10.1007/s11004-016-9658-6](https://doi.org/10.1007/s11004-016-9658-6)

Publication date

2017

Document Version

Final published version

Published in

Mathematical Geosciences

Citation (APA)

Wambeke, T., & Benndorf, J. (2017). A Simulation-Based Geostatistical Approach to Real-Time Reconciliation of the Grade Control Model. *Mathematical Geosciences*, 49(1), 1–37.
<https://doi.org/10.1007/s11004-016-9658-6>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

A Simulation-Based Geostatistical Approach to Real-Time Reconciliation of the Grade Control Model

T. Wambeke^{1,2}  · J. Benndorf³

Received: 21 April 2016 / Accepted: 24 September 2016 / Published online: 14 October 2016
© The Author(s) 2016. This article is published with open access at Springerlink.com

Abstract One of the main challenges of the mining industry is to ensure that produced tonnages and grades are aligned with targets derived from model-based expectations. Unexpected deviations, resulting from large uncertainties in the grade control model, often occur and strongly impact resource recovery and process efficiency. During operation, local predictions can be significantly improved when deviations are monitored and integrated back into the grade control model. This contribution introduces a novel realization-based approach to real-time updating of the grade control model by utilizing online data from a production monitoring network. An algorithm is presented that specifically deals with the problems of an operating mining environment. Due to the complexity of the material handling process, it is very challenging to formulate an analytical approximation linking each sensor observation to the grade control model. Instead, an application-specific forward simulator is built, translating grade control realizations into observation realizations. The algorithm utilizes a Kalman filter-based approach to link forward propagated realizations with real process observations to locally improve the grade control model. Differences in the scale of support are automatically dealt with. A literature review, following a detailed problem description, presents an overview of the most recent approaches to solving some of the practical problems identified. The most relevant techniques are integrated and the resulting mathematical framework is outlined. The principles behind the self-learning algorithm are explained. A synthetic experiment demonstrates that the algorithm is capable

✉ T. Wambeke
t.wambeke@tudelft.nl

¹ Delft University of Technology, Stevinweg 1, 2628CN Delft, The Netherlands

² IHC MTI, Delftechpark 13, 2628XJ Delft, The Netherlands

³ University of Technology Bergakademie Freiberg, Fuchsmühlenweg 9, 09599 Freiberg, Germany

of improving the grade control model based on inaccurate observations on blended material streams originating from two extraction points.

Keywords Grade control model · Reconciliation · Geostatistics · Ensemble Kalman filter

1 Introduction

Traditionally, the mining industry has had mixed successes in achieving the production targets it has set out. Several projects have been identified where mineral grades are not as expected, schedules and plans are not met and recovery is lower than forecasted (McCarthy 1999; Vallee 2000; Tatman 2001; McCarthy 2003). The deviations of the produced tonnages and grades from model-based expectations result from a mismatch between the scale of the exploration data and the short-term production targets (Bendorf 2013). In other words, it is challenging to accurately define the characteristics of, e.g., a few truckloads, designated to be transported to the processing plant, based on data gathered at relatively wide grids. For certain commodities, it is common to perform grade control (GC) drilling to further reduce the uncertainty (Peattie and Dimitrakopoulos 2013; Dimitrakopoulos and Godoy 2014). However, GC drilling is expensive and almost exclusively focused on sampling grades. Metallurgical properties are often ignored.

The mineral industry is increasingly looking for effective methods for monitoring and reconciling estimates and actual observations at different stages of the resource extraction process (Morley 2014). Recent developments in sensor technology enable the online characterization of production performance and raw material characteristics. To date, sensor measurements are mainly utilized in forward loops for downstream process control and material handling (Zimmer 2012; Lessard et al. 2014; Nienhaus et al. 2014). A backward integration of sensor information into the GC model to continuously improve the production forecasts and dispatch decisions does not yet occur.

The application of sensors carries a large potential regarding process improvements. Sensor responses could be used to progressively increase the knowledge about the in situ material characteristics. This has two main consequences. First, the frequency of misallocation could decrease, i.e., a smaller amount of actual ore is incorrectly allocated to the waste dump and a smaller amount of actual waste enters the processing plant. Second, an improved characterization of metallurgical properties could lead to a more optimal selection of process parameters. For example, the throughput of the comminution circuit can be reduced upfront when harder ore is expected to ensure that the resulting grain sizes stay within acceptable limits. A proactive selection of process parameters in combination with the elimination of low value material from the processing plant will further result in a reduction of dilution, an increase of concentrator recovery and a larger annual metal production.

The potential of real-time updating is obvious. To apply it in practice, algorithms need to be developed capable of assimilating direct and indirect measurements into the GC model. Thus, at any point in time when new observations become avail-

able, the following inverse problem needs to be solved (Tarantola 2005; Oliver et al. 2008)

$$\mathbf{z} = \mathcal{A}^{-1}(\mathbf{d}), \quad (1)$$

where $\mathcal{A}$ is a forward observation model (linear or nonlinear) that maps the spatial attributes $\mathbf{z}$ of the GC model onto sensor observations $\mathbf{d}$. The observations result from either direct or indirect measurements. The following challenges are identified: (1) the latest solution should account for previously integrated data (sequential approach); (2) due to the nature of a mining operation, it is nearly impossible to formulate an analytical approximation of the forward observation model, let alone compute its inverse and (3) observations are made on blended material streams originating from multiple extraction points. The objective of this paper is to present a new algorithm to assimilate sensor observations into the grade control model, specifically tailored to the requirements of the mining industry.

2 Literature Review

Kitanidis and Vomvoris (1983) introduced a geostatistical approach to the inverse problem in groundwater modeling. Both scarce direct (local log conductivity) and more abundant indirect (hydraulic head and arrival time) measurements are used to estimate the hydraulic conductivity in geological media through a linear estimation procedure known in the geostatistical literature as cokriging (Journel and Huijbregts 1978; Deutsch and Journel 1992). In the Bayesian literature, the same procedure is referred to as updating or conditioning (Schweppe 1973; Wilson et al. 1978; Dagan 1985). The geostatistical approach received considerable attention (Hoeksema and Kitanidis 1984; Rubin and Dagan 1987; Yates and Warrick 1987; Sun and Yeh 1992; Harter and Yeh 1996; Tong 1996).

Later, several simultaneous and independent developments resulted in a method to recursively incorporate subsets of data one at a time (Evensen 1992; Harvey and Gorelick 1995; Yeh and Zhang 1996). The proposed sequential estimator improves previous subsurface models by using linearly weighted sums of differences between observations and model-based predictions

$$\mathbf{z}_t = \mathbf{z}_{t-1} + \mathbf{W}_t(\mathbf{d}_t - \mathcal{A}_t(\mathbf{z}_{t-1})), \quad (2)$$

where the vector $\mathbf{z}_t$ contains estimates of the spatial attributes after t updates; the vectors $\mathbf{d}_t$ and $\mathcal{A}_t(\mathbf{z}_{t-1})$, respectively, hold actual observations and model-based predictions at time t and $\mathbf{W}_t$ is a matrix with kriging weights defining the contribution of the detected deviations to the updated subsurface model. If the vector $\mathbf{d}_t$ only contains direct local measurements, then the linear estimator corresponds to simple kriging, this is kriging with a known mean. On the other hand, if also indirect measurements are included, then a single update results from solving a system of cokriging equations (Goovaerts 1997; Chiles and Delfiner 2012). The sequential linear estimator bears a remarkable resemblance to Kalman filter techniques (Kalman 1960; Evensen

1992; Bertino et al. 2002). The kriging weights are computed from the forecast and observation error covariance matrices, $\mathbf{C}_{t-1, zd}$ and $\mathbf{C}_{t-1, dd}$

$$\mathbf{W}_t = \mathbf{C}_{t-1, zd} \mathbf{C}_{t-1, dd}^{-1} \quad (3a)$$

$$= \mathbf{C}_{t-1, zz} \mathbf{A}_t^T (\mathbf{A}_t \mathbf{C}_{t-1, zz}^{-1} \mathbf{A}_t^T + \mathbf{R})^{-1}, \quad (3b)$$

where $\mathbf{C}_{t-1, zz}$ is the prior error covariance matrix of the attribute field, $\mathbf{R}$ is a diagonal matrix which specifies the sensor precision (a large sensor precision corresponds to a low value on the diagonal) and $\mathbf{A}_t$ is a first-order approximation of the nonlinear observation model $\mathcal{A}_t$ (Evensen 1992; Yeh and Zhang 1996). If both the prior error covariance matrix and the sensor precision tend to be large, then the kriging weights tend to increase, indicating that a significant portion of the detected deviations are taken into account to update the attribute field. For completeness, the posterior error covariance matrix of the attribute field after one assimilation step is also given

$$\mathbf{C}_{t, zz} = (\mathbf{I} - \mathbf{W}_t \mathbf{A}_t) \mathbf{C}_{t-1, zz}. \quad (4)$$

Vargaz-Guzman and Yeh (1999) provided the theoretical evidence that under a linear observation model, sequential kriging and cokriging are equivalent to their traditional counterpart which includes all data simultaneously. In case of a nonlinear observation model, a sequential incorporation of data increases the accuracy of the first-order (linear) approximations $\mathbf{A}_t$, since they are calculated around a progressively improving attribute field (Evensen 1992; Harvey and Gorelick 1995). The linear estimator thus propagates the conditional mean and covariances from one update cycle to the next.

At time zero, a global covariance model suffices to describe the degree and scale of variability in the attribute field. The covariance matrix $\mathbf{C}_{(0, zz)}$ is stationary. At any other time, the updated covariances reflect the assimilation history and indirectly depend on the location of the material sources (Harvey and Gorelick 1995). This nonstationarity of the updated covariances results in perhaps the greatest limitation of the method, i.e., the necessity of storing large covariance matrices. Despite the promising results, the above-mentioned techniques were yet not considered for resource modeling and reconciliation.

Thus far, the discussion was focused on the sequential updating of a single best estimate. In geostatistics, it is common to simulate a set of realizations to assess uncertainty (Dowd 1994; Dimitrakopoulos 1998; Rendu 2002). The propagated conditional mean and covariance provide an intuitive description of statistics required to perform geostatistical simulations. For example, a single realization can be generated through the combination of the propagated mean with the product of a decomposed covariance matrix (LU decomposition) and a vector filled with white noise (Davis 1987; Alabert 1987). Gomez-Hernandez and Cassiraga (2000), Hansen et al. (2006) and Hansen and Mosegaard (2008) opt for a different approach and propose using the entire collection of measurements simultaneously in combination with a cokriging-based version of sequential Gaussian simulation. As time progresses, such an approach results in

significant memory usage due to the substantial growth of available production data. The simulation approaches discussed thus far all require that the simulation algorithm is completely rerun after each timestep.

Vargaz-Guzman and Dimitrakopoulos (2002) presented an approach that can facilitate fast updating of generated realizations based on new data, without repeating the full simulation process. The approach is termed conditional simulation of successive residuals and was designed to overcome the size limitations of the LU decomposition. The lower triangular matrix is obtained through a novel column partitioning, expressed in terms of successive conditional covariance matrices. The partitioning requires the specification of a sequence of (future) data locations. A stored L matrix can then facilitate the conditional updating of existing realizations if and only if the sequence of visited subsets and used production data are the same as the one used for generating the initial realizations (Jewbali and Dimitrakopoulos 2011; Dimitrakopoulos and Jewbali 2013).

A major limitation of the techniques discussed thus far results from the necessity to store and propagate the conditional nonstationary covariances. A significant portion of memory needs to be allocated to hold a collection of elements, the size of the square of the number of grid nodes. The very large grids commonly encountered in the mining industry make such approaches infeasible. To circumvent these limitations, the previously discussed sequential linear estimator (Eq. 2) can be integrated into a Monte Carlo framework (Evensen 1994). At each time step, a finite set of realizations is updated on the basis of sensor observations collected at time t :

$$\mathbf{Z}_t(:, i) = \mathbf{Z}_{t-1}(:, i) + \mathbf{W}_t(\mathbf{d}_t - \mathcal{A}_t(\mathbf{Z}_{t-1}(:, i)) + \mathbf{E}_t(:, i)) \quad \forall i \in I, \quad (5)$$

where the observation errors $\mathbf{E}_t(:, i)$ are randomly drawn from a normal distribution with a zero mean vector and a diagonal covariance matrix $\mathbf{R}$. The conditional forecast and observation error covariances, $\mathbf{C}_{t-1,zd}$ and $\mathbf{C}_{t-1,dd}$, are computed empirically from both sets of model-based predictions $\mathcal{A}_t(\mathbf{Z}_{t-1}(:, i))$ and field realizations $\mathbf{Z}_{t-1}(:, i)$. Due to the applied Monte Carlo concept, the first-order approximation of the forward observation model can be avoided (Evensen 1997; Burgers and van Leeuwen 1998). The initial set of realizations can be generated using techniques of conditional simulation. Aside from the empirical calculation of the covariances, the Monte Carlo-based sequential conditioning approach bears some resemblance to the equations for conditioning simulations as presented in Journel and Huijbregts (1978). Applications of the sequential conditioning approach in a geoscientific context can be found in many documented studies (Jansen et al. 2012; Heidari et al. 2011; Franssen et al. 2011; Hu et al. 2012; Bertino et al. 2002). These other research disciplines refer to this technique as the ensemble Kalman filter.

Before proceeding, it is important to comprehend the subtle differences to previous applications of the (ensemble) Kalman filter. Weather forecasting and reservoir modeling (oil, gas and water) consider dynamic systems repetitively sampled at the same locations. Generally, each observation characterizes a volume surrounding a sample location. These local volumes are sampled repetitively in time. Mineral resource modeling on the other hand focuses on static systems gradually sampled at different locations. Each observation is characteristic for a blend of material originating from

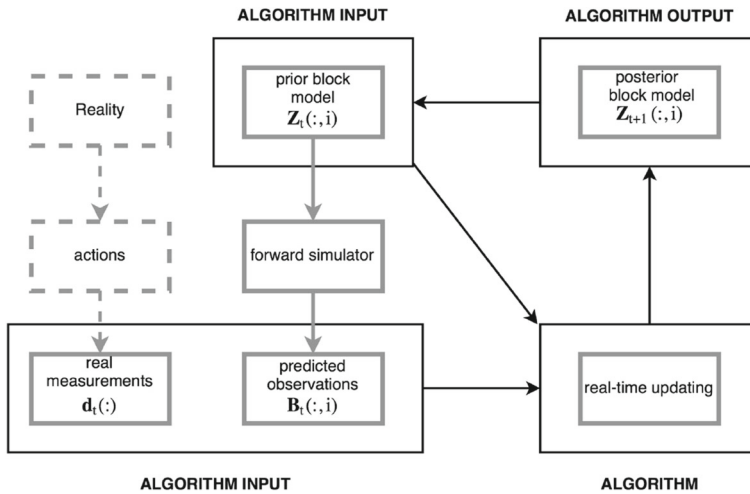


Fig. 1 Overview of the interaction between the updating algorithm and the forward simulator

one or more extraction points. Each part of the material stream is sampled only once, at the moment it passes the sensor in the production chain. Although virtually the same, the terminology of sequential updating is preferred over (ensemble) Kalman filtering. This is to stress the absence of dynamic components and to highlight its link to the field of geostatistics.

Benndorf (2015) was the first to recognize the potential of sequential updating in a context of mineral resource extraction. His preliminary work provides a mathematical description linked to the unique configurations of the mining industry (as discussed above). However, the practicality of the description is limited due to the need for an explicit formulation of the A_t matrix. The presented description is based on 3b, which requires an empirical computation of the full $C_{t-1,zz}$ matrix together with its inverse. Hence, the description results in high computational costs and memory requirements. Moreover, the diagonal R matrix characterizing the sensor precision needs to be formulated in the computation domain: straightforward when all distributions involved are Gaussian, much less intuitive when transformations are applied.

This contribution presents an algorithm based on Eq. 3a. The forward observation model $\mathcal{A}$ is excluded from the computer code (Fig. 1). Instead, a forward simulator is used to convert field realizations into model-based predictions ($\mathcal{A}$ is applied to fields individually rather than to the covariance matrix). The use of a forward simulator overcomes the challenges of formulating an analytical approximation linking each sensor observation to the grade control model. The necessary forecast and observation error covariances, $C_{t-1,dz}$ and $C_{t-1,dd}$, can now be computed empirically. Differences in the scale of support are inherently accounted for. A Gaussian anamorphosis option is implemented to deal with suboptimal conditions related to indirect observations and non-Gaussian distributions. A specific algorithm structure ensures that the measurement error can be defined on its original units and does not need to be translated into a normal score equivalent. An interconnected parallel updating sequence (helix) can

be configured to reduce the effects of filter inbreeding (covariance collapse). A neighborhood option can be activated to further reduce computation times and memory requirements. Two covariance correction strategies are implemented to contain the propagation of statistical sampling errors originating from the empirical computation of covariances.

3 Algorithm

This section further elaborates on the design of the algorithm and its specific options. The first subsection discusses the rationale behind the Gaussian anamorphosis. The implementation, allowing for an intuitive treatment of measurement error, is explained. The second subsection provides a computationally efficient strategy for solving the updating equations. The presentation is generic, in that it includes the option to configure a parallel updating sequence (helix). The third subsection discusses the neighborhood option. The fourth subsection provides some insight into both covariance correction techniques. The fifth subsection presents the pseudocode of the updating algorithm. The pseudocode illustrates how the individual functional components are integrated.

3.1 Gaussian anamorphosis

The updating approach presented in Benndorf (2015) is only optimal if all involved variables are multivariate Gaussian and if the forward observation model $\mathcal{A}$ is linear. Commonly, global normal score transformations are used to mitigate the effects of multivariate distributions which are not Gaussian (Goovaerts 1997; Deutsch and Journel 1992). However due to the developing non-stationarity, the underlying assumptions get violated. This is, after a few updates, the global cumulative field distribution does not anymore represent the local updated conditions. Instead, grid nodes are individually transformed according to node-specific Gaussian anamorphosis functions (Beal et al. 2010; Simon and Bertino 2009). This approximation just renders marginal distributions univariate Gaussian. Multivariate distributions are not explicitly corrected. Yet, the random vectors are closer to being multivariate Gaussian than prior to the univariate transformations (Zhou et al. 2011). First, the univariate transformation of grid nodes is discussed. The transformation of observations and the treatment of measurement error are addressed later on.

At each grid node $\mathbf{x}_n$, a local Gaussian anamorphosis function $\Phi_{\mathbf{x}_n}$ is built and used to forward transform the Monte Carlo representation of the random vector $\mathbf{z}(\mathbf{x}_n)$

$$\mathbf{U}_t(n, i) = \Phi_{t, \mathbf{x}_n}(\mathbf{Z}_t(n, i)) \quad (6a)$$

$$= G^{-1}(F(\mathbf{Z}_t(n, i))) \quad \forall i \in I, \quad (6b)$$

where I refers to the total number of elements in the Monte Carlo sample. A two-step procedure is used for the implementation of the forward transformation. (1) The cumulative probability $F(\mathbf{Z}_t(n, i))$ is estimated based on the rank i' of the corresponding

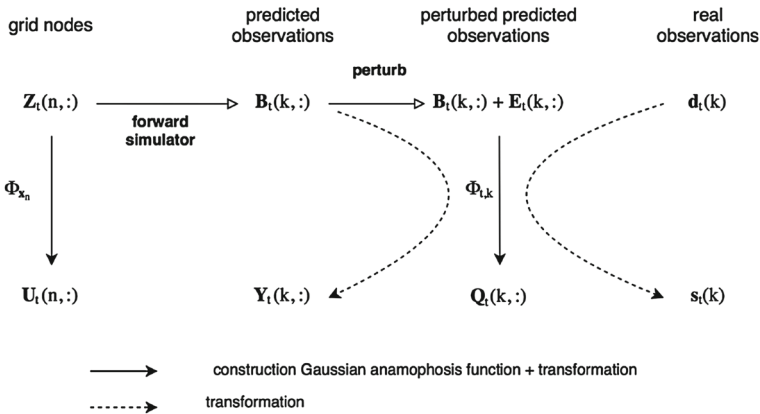


Fig. 2 Univariate transformation of grid nodes, perturbed predicted observations, predicted observations and actual observations

element in the full Monte Carlo sample ($\frac{i'-0.5}{I}$). (2) The resulting cumulative probability is mapped onto its corresponding standard normal score value. In other words, the inverse of the standard normal cumulative distribution function G is computed. The left hand side of Fig. 2 shows the procedure for the univariate transformation of grid nodes.

The necessity and added value of a transformation of observations depends partially on the considered updating scheme (Eqs. 3a vs. 3b). To further clarify this statement, a situation with transformed grid nodes is assumed and the impact of non-Gaussian observations on both updating schemes is analyzed.

A study of Eq. 3b shows that the kriging weights W_t can be computed without explicitly considering the predicted observations. The transformed grid nodes $U(n, :)$ are used to compute the empirical field error covariances $C_{t-1,uu}$. Both forecast error and observation error covariances are then written as linear combinations hereof. The resulting matrices, $C_{t-1,uu}A_t^T$ and $A_tC_{t-1,uu}^{-1}A_t^T + R$ are fully described in a normal score space. As a consequence, sensor precisions need to be characterized in the normal score domain as well (R matrix). In conclusion, non-Gaussian observations thus have a limited impact on the computation of the weights W_t when using updating Scheme 3b.

In contrast, due to the specific formulation of updating Scheme 3a, predicted observations and their distributions do have a direct and significant impact on the computation of the weights W_t . The forecast error and observation error covariances are after all empirically calculated based on untransformed predicted observations. They are no longer expressed as linear combinations of field error covariances, computed from transformed grid nodes. The use of a flexible forward simulator thus results in a less far-reaching effect of transformed grid nodes when computing the weights W_t .

Since updating Scheme 3a is less robust against non-Gaussian observations, a transformation is warranted. This will further increase the performance of the sequential updating technique inasmuch as all involved variables, including the observations,

better approximate a multivariate Gaussian distribution. Moreover, a Gaussian anamorphosis of the observations leads to an apparent pseudolinearization of the forward observation model in the normal score domain (Schoniger et al. 2012). This pseudolinearization can be more accurately exploited by the linear updating equations.

The K predicted observations at time step t are defined as

$$\mathbf{B}_t(:, i) = \mathcal{A}_t(\mathbf{Z}_{t-1}(:, i)) \quad \forall i \in I, \quad (7)$$

where K refers to the number of measurements taken over a period $[t-1, t]$ and $\mathcal{A}_t$ represents a forward simulator translating single field realizations $\mathbf{Z}_t(:, i)$ into vectors $\mathbf{B}_t(:, i)$ with K predicted observations. To account for measurement error, the predicted observations are randomly perturbed, that is, I vectors $\mathbf{E}_t(:, i)$ drawn from $N(0, \mathbf{R})$ are added to the columns $\mathbf{B}_t(:, i)$. Note that the diagonal $\mathbf{R}$ matrix directly describes the measurement error in the original domain, and a translation to a normal score space is not needed. The rows of the resulting perturbed predicted observation matrix are transformed individually by evaluating K Gaussian anamorphosis functions (Fig. 2)

$$\mathbf{Q}_t(k, i) = \Phi_{t,k}(\mathbf{B}_t(k, i) + \mathbf{E}_t(k, i)) \quad \forall i \in I. \quad (8)$$

The addition of random noise $\mathbf{E}_t(:, i)$ results in a kernel smoothing effect making the estimation of the empirical CDF more robust, even for small sample sizes (Cheng and Parzen 1997). Each of the K actual inaccurate observations are assumed to originate from a distribution described by its corresponding Monte Carlo sample of perturbed predicted observations $(\mathbf{B}_t(k, :) + \mathbf{E}_t(k, :))$. Hence, the actual observations are transformed according to the previously constructed anamorphosis functions (Fig. 2)

$$\mathbf{s}_t(k) = \Phi_{t,k}(\mathbf{d}_t(k)). \quad (9)$$

The integration of measurement noise into the transformations allows for a smooth computation of covariances and deviations, without the need to further consider measurement error in the normal score domain. The normal score equivalent $\mathbf{R}$ matrix is implicitly accounted for when computing observation error covariances from the rows of the $\mathbf{Q}_t$ matrix. Previously, a noise vector had to be explicitly added to the deviations between the predicted and actual observations (Eq. 5). Now, a normal score equivalent is automatically inserted into the differences between $\mathbf{s}_t$ and the columns $\mathbf{Q}_t(:, i)$.

The currently available transformed data would not allow for a correct empirical computation of the forecast error covariances. A pairwise consideration of rows from the $\mathbf{U}$ and $\mathbf{Q}$ matrix would result in overestimated forecast error covariances. This is because the forecast error covariances should be computed without the consideration of measurement noise. Therefore, the previously constructed anamorphosis functions $\Phi_{t,k}$ are used once more to transform unperturbed predicted observations (Fig. 2)

$$\mathbf{Y}_t(k, i) = \Phi_{t,k}(\mathbf{B}_t(k, i)) \quad \text{for } i = 1, 2, \dots, I. \quad (10)$$

After the updating step, the grid nodes are backtransformed using the inverse of the grid node-specific anamorphosis functions Φ_{t,x_n}^{-1} . Back transformation of the observations is not required. Instead, updated predicted observations are obtained from a subsequent run of the forward simulator.

3.2 Solving the updating equations

The following provides an efficient strategy for solving the updating equations. The presentation is generic in that it includes an option to configure a pair of connected updating cycles. This option can be activated to avoid problems of inbreeding. Inbreeding is essentially caused by the empirical estimation of covariances from Monte Carlo samples of limited size (Zhou et al. 2011). When occurring, the spread in the Monte Carlo sample becomes unrealistically small; in particular the spread would be smaller than the difference between the sample mean and the unknown truth (Zhou et al. 2011). The problem generally arises when a group of grid nodes is being updated frequently during a certain period of time.

The idea of configuring a pair of sequential updating cycles originates from Houtekamer and Mitchell (1998), who argued that the computation of the weights $\mathbf{W}_t$ and the evaluation of their quality (diagonal elements in $\mathbf{C}_{t+1,uu}$) would be based on exactly the same information when using a single Monte Carlo sample. As the assimilation cycle proceeds, the propagation of this dependent quality test might eventually result in a collapse of the Monte Carlo sample. To maintain sufficient spread, two connected updating cycles can be configured. The weights $\mathbf{W}_t$ computed from one cycle are used to assimilate data into the other cycle. In this way, each of the two cycles use different Monte Carlo samples to estimate the weights and to evaluate the quality of an update.

To configure two connected updating cycles, the previously transformed Monte Carlo samples need to be split into two subsets. The constructed matrices with I columns ($\mathbf{U}$, $\mathbf{Q}$ and $\mathbf{Y}$) are subdivided into two new matrices with, respectively, α and $I - \alpha$ columns. Considering an original matrix $\mathbf{X}$, the new matrices are denoted by $\mathbf{X}(:, : \alpha)$ and $\mathbf{X}(:, \alpha :)$. The operators $: \alpha$ and $\alpha :$ refer, respectively, to the first α and the last $I - \alpha$ columns. The default value of the integer parameter α is set to $I/2$. Note that the two subsets do not necessarily have to be of equal size.

Once the Monte Carlo samples are split, all necessary covariances are empirically calculated. The first two $K \times K$ matrices with observation error covariances are constructed

$$\mathbf{C}_{t,q}^1(k, k) = E[(\mathbf{Q}_t(k, : \alpha) - E[\mathbf{Q}_t(k, : \alpha)])(\mathbf{Q}_t(k, : \alpha) - E[\mathbf{Q}_t(k, : \alpha)])], \quad (11a)$$

$$\mathbf{C}_{t,q}^2(k, k) = E[(\mathbf{Q}_t(k, \alpha :) - E[\mathbf{Q}_t(k, \alpha :)])(\mathbf{Q}_t(k, \alpha :) - E[\mathbf{Q}_t(k, \alpha :)])]. \quad (11b)$$

The observation error covariances are computed from transformed perturbed predicted observations and as such implicitly account for a normal score equivalent of the diagonal R matrix. Likewise, two $N \times K$ matrices with forecast error covariances are constructed:

$$\mathbf{C}_{t,uy}^1(n, k) = E[(\mathbf{U}_t(n, : \alpha) - E[\mathbf{U}_t(n, : \alpha)])(\mathbf{Y}_t(k, : \alpha) - E[\mathbf{Y}_t(k, : \alpha)])], \quad (12a)$$

$$\mathbf{C}_{t,uy}^2(n, k) = E[(\mathbf{U}_t(n, \alpha :) - E[\mathbf{U}_t(n, \alpha :)])(\mathbf{Y}_t(k, \alpha :) - E[\mathbf{Y}_t(k, \alpha :)])]. \quad (12b)$$

It is important to stress that the transformation of the unperturbed predicted observations (Eq. 10) is only performed to calculate the forecast error covariances.

A significant speedup of the updating operation can be achieved by solving the equations first in the K -dimensional ‘observation space’ and then projecting the solutions into the N -dimensional ‘field space’ (Evensen and van Leeuwen 1996; Cohn et al. 1998). A Cholesky decomposition of the two observation error covariance matrices results in considerable computation savings when solving the following equations

$$\mathbf{C}_{t,qq}^2 \mathbf{T}(:, i') = \mathbf{s}_t - \mathbf{Q}_t(:, i') \quad \forall i' \in [1, \alpha], \quad (13a)$$

$$\mathbf{C}_{t,qq}^1 \mathbf{T}(:, i'') = \mathbf{s}_t - \mathbf{Q}_t(:, i'') \quad \forall i'' \in [\alpha, I]. \quad (13b)$$

The two resulting lower triangular matrices are combined with deviations between transformed actual and predicted observations of the complementary subset to calculate a series of solution vectors $\mathbf{T}(:, i)$. Both series of solution vectors are subsequently combined with previously calculated forecast error covariances to update existing realizations:

$$\mathbf{U}_{t+1}(:, i') = \mathbf{U}_t(:, i') + \mathbf{C}_{t,uy}^2 \mathbf{T}(:, i') \quad \forall i \in [1, \alpha], \quad (14a)$$

$$\mathbf{U}_{t+1}(:, i'') = \mathbf{U}_t(:, i'') + \mathbf{C}_{t,uy}^1 \mathbf{T}(:, i'') \quad \forall i'' \in [\alpha, I]. \quad (14b)$$

Finally, the Monte Carlo samples of grid nodes are recombined and transformed back using the grid node-specific anamorphosis functions which were stored in memory

$$\mathbf{Z}_{t+1}(n, i) = \Phi_{t,x_n}^{-1}(\mathbf{U}(n, i)) \quad \forall i \in I. \quad (15)$$

3.3 Neighborhood

Often, it is impractical (especially for large grids) to consider all N grid nodes simultaneously for updating, when K observations become available. Since the spatial correlation generally decreases with distance, it is sufficient to only include a node when it is located in the near vicinity of one of the L extraction points. A search algorithm based on a space-partitioning tree structure is implemented to allocate grid nodes to a neighborhood set N' replacing the full grid N in all subsequent computations. A local neighborhood approach is not new. However, the current implementation requires that neighborhoods around active extraction points are considered simultaneously. For example, due to measurements on blended material streams, an observation can be correlated to multiple nodes from different neighborhoods.

The subset N' is constructed prior to performing the transformations. As a result, the order of the $\mathbf{U}$ matrices in Eqs. 6, 12, 14 and 15 greatly reduces permitting a substantial computational economy to be realized. A local neighborhood strategy also

reduces the effects of spurious correlations which may arise between observations and remote grid nodes. Eventually, spurious correlations might lead to an update of several remote nodes based on irrelevant information. To avoid such a fallacy, a neighborhood strategy can be used to exclude remote nodes from the analysis.

3.4 Covariance error correction

The previous subsections already touched upon some of the issues related to the empirical estimation of covariances from finite Monte Carlo samples. A standard result from statistical theory states that the correlation between two normal distributions can be estimated with the following accuracy when using I Monte Carlo pairs

$$E[(\rho - \hat{\rho})^2] \approx \frac{(1 - \hat{\rho})^2}{I}, \quad (16)$$

where ρ and $\hat{\rho}$ are, respectively, the unknown and estimated correlation. The term $E[(\rho - \hat{\rho})^2]$ expresses the variance around the estimate. It follows that an accurate estimation of weak correlations (or covariances) would in any case require very large sample with thousands of realizations. When due to practical considerations the Monte Carlo samples are rather limited, numerical inaccuracies may arise. To mitigate possible numerical inaccuracies, two covariance correction techniques are implemented: an error correction and a localization correction. Both techniques apply a differential correction to the individual elements of the forecast error covariance matrix $\mathbf{C}_{uy}$. The magnitude of the correction (indirectly) depends on the estimated value of each covariance element.

The localization correction is based on the notion that spatial covariance generally decreases with distance. Equation 16 indicates that the accuracy of a covariance estimate reduces accordingly. To account for distant dependent numerical inaccuracies, the forecast error covariance matrix $\mathbf{C}_{t,uy}$ is localized. The concept of localization requires observations which can be attributed to a certain spatial location. As a result, a physically sensible distance between an observation k and a grid block n can be calculated (Sakov and Betino 2011). Given a taper function ρ , a correction factor $\rho(n, k)$ is computed and used to correct the corresponding covariance ($\mathbf{C}_{t,uy}(n, k) * \rho(n, k)$). Since in mining applications, observations are made on blended material streams, the conventional implementation had to be adjusted. The following assumes that the locations of the L simultaneous extraction points are known. The localized element of the $\mathbf{C}_{uy}$ matrix can then be written as a sum of the product between grid realization anomalies ($\mathbf{U}_t(n, :) - E[\mathbf{U}_t(n, :)]$) tapered around extraction location l and observation anomalies ($\mathbf{Y}_t(k, :) - E[\mathbf{Y}_t(k, :)]$)

$$\mathbf{C}_{t,uy}(n, k) = \sum_{l=1}^L \left(E \left[\rho(n, l) * (\mathbf{U}_t(n, :) - E[\mathbf{U}_t(n, :)]) * \right. \right. \\ \left. \left. (\mathbf{Y}_t(k, :) - E[\mathbf{Y}_t(k, :)]) \right] \right). \quad (17)$$

The Gaspari–Cohn correlation function is used for tapering the grid realization anomalies (Gaspari and Cohn 1999). Its use requires the definition of three range parameters (x , y , z direction). A good choice for the range parameters is based on the number of realizations and the prior covariance (Chen and Oliver 2010). As the number of realizations increases, the accuracy of covariance estimates increases and thus larger range parameters can be selected.

The error correction technique uses an off-line Monte Carlo simulation to characterize the potential variations between unknown and estimated correlations (Eq. 16). A lookup table is constructed linking the estimated correlations with the correction factors between 0 and 1 Anderson (2012). The use of a lookup table requires a fixed value range $[-1, 1]$; hence covariances are replaced by correlations without any further consequences. The procedure for constructing the lookup table is as follows. For every discrete correlation ρ_m in the range $[-1, 1]$, a Monte Carlo experiment is performed to compute the corresponding correction factor. Each Monte Carlo experiment consists of five steps. (i) Define a bivariate normal distribution with zero mean and covariance $[[1, \rho_m], [\rho_m, 1]]$. (ii) Draw I times two values from this distribution. Compute the correlation coefficient based on the I pairs. (iii) Repeat (ii) P times. (iv) Compute the mean $\hat{\rho}_m$ and standard deviation σ_m of the P estimated correlation values. (v) Define α as $(\hat{\rho}_m/\sigma_m)^2$. The correction factor (CF) for an estimated correlation of ρ_m is then computed as

$$CF(\rho_m) = \frac{\alpha}{1 + \alpha} * \frac{\hat{\rho}_m}{\rho_m}. \quad (18)$$

The described Monte Carlo experiment needs to be repeated M times to construct the lookup table. The parameters M and P are, respectively, set to 201 and 10^6 . I should equal the number of realizations used to estimate the empirical correlations.

Once a lookup table is constructed or read from the file, the error correction can be applied to the $\mathbf{C}_{t,uy}$ matrix

$$\mathbf{C}_{t,uy}(n, k) = CF\left(\frac{\mathbf{C}_{t,uy}(n, k)}{\mathbf{S}_{t,uu}(n, n)\mathbf{S}_{t,yy}(k, k)}\right) * \mathbf{C}_{t,uy}(n, k). \quad (19)$$

This requires the additional computation of the empirical standard deviation at N grid nodes ($\mathbf{S}_{t,uu}(n, n)$) and K observations ($\mathbf{S}_{t,yy}(k, k)$).

3.5 Pseudocode

The following pseudocode illustrates how the previously discussed functional components are integrated in one updating procedure (Algorithm 1). The option of using two interconnected updating cycles has been omitted to avoid unnecessary complexity. To accommodate for this option, lines 26 to 44 (Algorithm 1) have to adjusted based on Eqs. 11 to 14.

Algorithm 1 Updating Algorithm

```

1: procedure UPDATE( $\mathbf{Z}_t, \mathbf{d}_t, \mathbf{B}_t, \Sigma$ ) ▷  $\Sigma$  holds  $L$  extraction locations  $\mathbf{x}_l$ 
2:    $tree \leftarrow \text{CONSTRUCTTREE}(\mathbf{Z}_t)$ 
3:   for  $\mathbf{x}_l$  in  $\Sigma$  do ▷ collect nodes around  $\mathbf{x}_l$ 
4:      $\mathbf{Z}_t^l \leftarrow \text{SEARCHTREE}(tree, \mathbf{x}_l)$  ▷  $\mathbf{Z}_t^l$  is a subset of  $\mathbf{Z}_t$ 
5:      $\mathbf{Z}_t^\Sigma \leftarrow \text{UNION}(\mathbf{Z}_t^\Sigma, \mathbf{Z}_t^l)$  ▷ combine subsets

6:    $\mathbf{U}_t \leftarrow \text{INITIALIZE}(N', I)$  ▷  $N'$  nodes in combined subsets
7:   for  $n$  in  $N'$  do
8:      $\Phi_{t, \mathbf{x}_n} \leftarrow \text{CONSTRUCT}(\mathbf{Z}_t^\Sigma(n, :))$ 
9:     for  $i$  in  $I$  do
10:       $\mathbf{U}_t(n, i) \leftarrow \Phi_{t, \mathbf{x}_n}(\mathbf{Z}_t^\Sigma(n, i))$ 

11:   $\mathbf{Q}_t \leftarrow \text{INITIALIZE}(K, I)$  ▷  $K$  observations
12:   $\mathbf{Y}_t \leftarrow \text{INITIALIZE}(K, I)$ 
13:   $\mathbf{s}_t \leftarrow \text{INITIALIZE}(K, 1)$ 
14:   $\mathbf{E}_t \leftarrow \text{GENERATE}(\text{sensor accuracy})$  ▷  $K \times I$  matrix with white noise
15:   $\mathbf{F}_t \leftarrow \mathbf{B}_t + \mathbf{E}_t$ 
16:  for  $k$  in  $K$  do
17:     $\Phi_{t,k} \leftarrow \text{CONSTRUCT}(\mathbf{F}_t(k, :))$ 
18:    for  $i$  in  $I$  do
19:       $\mathbf{Q}_t(k, i) \leftarrow \Phi_{t,k}(\mathbf{F}_t(k, i))$  ▷ transform perturbed predictions
20:       $\mathbf{Y}_t(k, i) \leftarrow \Phi_{t,k}(\mathbf{B}_t(k, i))$  ▷ transform predictions
21:       $\mathbf{s}_t(k) \leftarrow \Phi_{t,k}(\mathbf{d}_t(k))$  ▷ transform real observations

22:   $\mathbf{C}_{t,qq} \leftarrow \text{INITIALIZE}(K, K)$ 
23:  for  $k_i$  in  $K$  do
24:    for  $k_j$  in  $K$  do
25:       $\mathbf{C}_{t,qq}(k_i, k_j) \leftarrow \text{COV}(\mathbf{Q}_{t,qq}(k_i, :), \mathbf{Q}_{t,qq}(k_j, :))$  ▷ Eq. 11

26:   $\mathbf{C}_{t,uy} \leftarrow \text{INITIALIZE}(N', K)$ 
27:  for  $n$  in  $N'$  do
28:    for  $k$  in  $K$  do
29:      if CovarianceLocalization is FALSE then
30:         $\mathbf{C}_{t,uy}(n, k) \leftarrow \text{COV}(\mathbf{U}_t(n, :), \mathbf{Y}_t(k, :))$  ▷ Eq. 12
31:      else if CovarianceLocalization is TRUE then
32:         $\mathbf{C}_{t,uy}(n, k) \leftarrow \text{COVLOC}(\mathbf{U}_t(n, :), \mathbf{Y}_t(k, :), \Sigma)$  ▷ Eq. 17

33:  if ErrorCorrection is TRUE then
34:    for  $n$  in  $N'$  do
35:       $\mathbf{S}_{t,uu}(n, n) \leftarrow \text{STANDARDDEVIATION}(\mathbf{U}_t(n, :))$ 
36:      for  $k$  in  $K$  do
37:         $\mathbf{S}_{t,yy}(k, k) \leftarrow \text{STANDARDDEVIATION}(\mathbf{Y}_t(k, :))$ 
38:         $\rho \leftarrow \mathbf{C}_{t,uy}(n, k) / (\mathbf{S}_{t,uu}(n, n) * \mathbf{S}_{t,yy}(k, k))$ 
39:         $CF \leftarrow \text{LOOPUP}(\rho)$  ▷ from table
40:         $\mathbf{C}_{t,uy}(n, k) \leftarrow CF * \mathbf{C}_{t,uy}(n, k)$  ▷ correct & replace, Eq. 19

41:   $\mathbf{T} \leftarrow \text{INITIALIZE}(K, I)$ 
42:  for  $i$  in  $I$  do
43:     $\mathbf{T}(:, i) \leftarrow \text{SOLVE}(\mathbf{C}_{t,qq} \mathbf{T}(:, i) = \mathbf{s}_t - \mathbf{Q}_t(:, i))$ 
44:     $\mathbf{U}_{t+1}(:, i) \leftarrow \mathbf{U}_t(:, i) + \mathbf{C}_{t,uy} \mathbf{T}(:, i)$ 

45:  for  $n$  in  $N'$  do
46:    for  $i$  in  $I$  do
47:       $\mathbf{Z}_{t+1}^\Sigma(n, i) \leftarrow \Phi_{t, \mathbf{x}_n}^{-1}(\mathbf{U}_{t+1}(n, i))$ 
48:   $\mathbf{Z}_{t+1} \leftarrow \text{OVERWRITE}(\mathbf{Z}_{t+1}^\Sigma, \mathbf{Z}_t)$  ▷ Overwrite the  $N'$  selected nodes
49:  Return  $\mathbf{Z}_{t+1}$ 

```

4 Experiment

4.1 General Setup

A synthetic experimental environment is created to showcase the performance of the developed algorithm. The synthetic environment represents a mining operation with two extraction points of unequal production rate. The resulting material streams are blended and inaccurate observations are made. The goal of the experiment is to test whether the algorithm is capable of improving the GC model based on unequally blended inaccurate observations. To stress the synthetic nature of the experiment, units for the modeled attribute are omitted.

4.2 Reference Field

During the course of the artificial experiment, a true but unknown state is sampled twice. First, small point samples are collected on a regular grid prior to the construction of the GC model (exploration phase). Second, observations are made on a blend of multiple mining blocks (operational phase). To correctly account for both scales of support, two representations of the true but unknown field are constructed. Both the high-resolution point and the lower-resolution block reference field characterize the true state over an area of $300 \text{ m} \times 300 \text{ m}$ using a different level of discretization.

The high-resolution point reference field is constructed on a discretized grid with 300×300 cells of size $1 \text{ m} \times 1 \text{ m}$. This reference field is generated using a random field simulation algorithm (sequential Gaussian simulation). The reference field is simulated based on an isotropic exponential function with a variance of 1.0 and a range of 100 m. Due to field dimensions which are relatively small compared to the range of the covariance model (300 m vs. 100 m), a four times larger simulation grid is used (600×600 cells of size $1 \text{ m} \times 1 \text{ m}$). Eventually, only the south-west quarter of the grid is retained (300×300 cells of size $1 \text{ m} \times 1 \text{ m}$). The resulting field has an average of -0.251 (-0.052) and a variance of 0.862 (1.006). The 5 and 95 % quantiles are, respectively, -1.733 (-1.733) and 1.312 (1.518). The values between brackets refer to those computed from the larger simulation grid.

The low-resolution block reference field is constructed on a discretized grid with 60×60 cells of size $5 \text{ m} \times 5 \text{ m}$. This second reference field is obtained from reblocking the previous point field. Values of 5×5 , $1 \text{ m} \times 1 \text{ m}$ cells are averaged and assigned to larger $5 \text{ m} \times 5 \text{ m}$ blocks. The resulting block reference field has an average of -0.251 and a variance of 0.792 . The 5 and 95 % quantiles are, respectively, -1.696 and 1.246 . As expected, the variability reduces due to reblocking. The final block reference field is displayed in Fig. 3a.

4.3 Prior Set of Realizations

A set of field realizations is used throughout to characterize the spatial variability and quantify the geological uncertainty of the study environment. The geological uncertainty originates from the limited amount of data available. Prior to the operation,

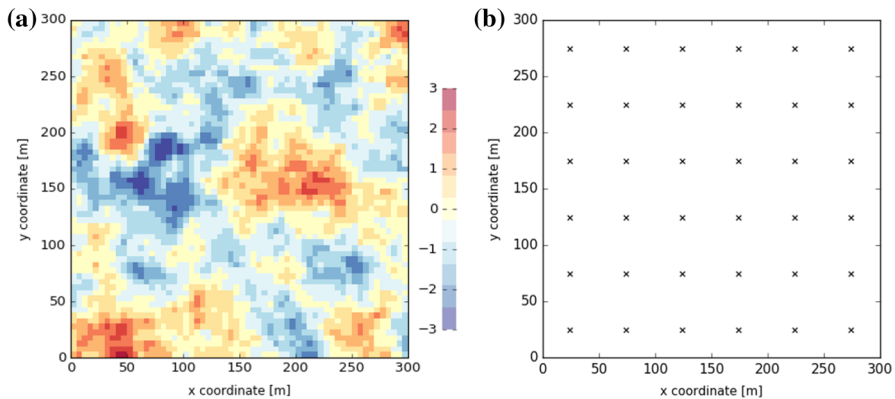


Fig. 3 True but unknown field **a** block representation, 60×60 blocks of size $5 \text{ m} \times 5 \text{ m}$ color coded according to their average block value, **b** location of point samples on $50 \text{ m} \times 50 \text{ m}$ exploration grid-point samples have a support of $1 \text{ m} \times 1 \text{ m}$

the study prescribes the availability of exploration data collected on a $50 \text{ m} \times 50 \text{ m}$ sampling grid. This medium-dense sampling grid (spacing equals half the range of the covariance model) results in a prior set of realizations with a realistic level of uncertainty and error.

A prior set of 200 realizations is simulated on a discretized grid with 600×600 $1 \text{ m} \times 1 \text{ m}$ cells. Again, the simulation grid is four times larger than the studied area for reasons of small field dimensions compared to the range of the covariance model. A total of 144 sampling points are provided as conditioning data (sampled on a $50 \text{ m} \times 50 \text{ m}$ exploration grid from the previous point simulation on the large 600×600 grid with $1 \text{ m} \times 1 \text{ m}$ cells). From the 144, only 36 sampling points are actually located inside the study area (Fig. 3b). The covariance model of the true field is assumed to be known and was not inferred from the sampling points (isotropic exponential function with a variance of 1.0 and a range of 100 m). Once simulated, only the south-west quarter of the grid is retained (300×300 cells of size $1 \text{ m} \times 1 \text{ m}$). Subsequently, the 200 point simulations are reblocked onto a grid with 60×60 blocks of size $5 \text{ m} \times 5 \text{ m}$. Figure 4 displays four arbitrarily chosen block realizations.

Only the block realizations on a discretized grid with 60×60 blocks of size $5 \text{ m} \times 5 \text{ m}$ are further considered (one reference field $\mathbf{Z}^*$ and 200 prior block realizations $\mathbf{Z}_0(:, i)$). All previous point simulations were just an aid to this end.

Figure 5a and b display a mean and standard deviation (SD) field computed from the 200 prior realizations. The mean field, representing the current best block estimate, is considerably smoother than the reference field (compare Figs. 5a, 3a). The general larger-scale patterns are nevertheless reasonably reproduced. This is due to the availability of sufficient conditioning data (on a point support). The approximate locations of the sampling points are marked by a lower block standard deviation (compare Figs. 3b, 5b). Due to the effects of differences in the scale of support, the SD of $5 \text{ m} \times 5 \text{ m}$ blocks does not entirely reduce to zero when conditioning based on cells of size $1 \text{ m} \times 1 \text{ m}$. The minimum and maximum observed SD values amount to 0.283 and 0.907. Figure 5c displays the absolute difference between true but unknown and esti-

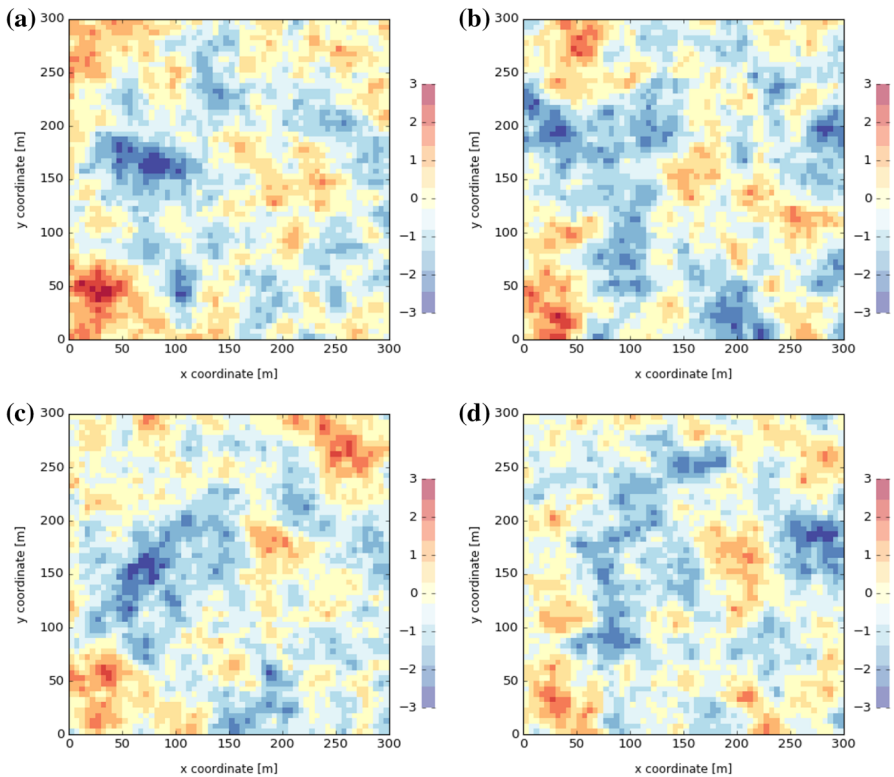


Fig. 4 Four arbitrarily chosen block realizations (no. 39, 56, 125 and 183), 60×60 blocks of size $5 \text{ m} \times 5 \text{ m}$ color coded according to their average block value

mated block values. The absolute error (AE) field is nothing more than the absolute difference between Figs. 3a and 5a. The figure illustrates that large deviations do occur locally, despite the considerable exploration effort. For example, the absolute error is consistently larger than 2.0 in an area bounded by $[25 \text{ m}, 175 \text{ m}] \times [75 \text{ m}, 225 \text{ m}]$.

Figure 6a, b depicts, respectively, a north–south ($x = 57.5 \text{ m}$) and east–west section ($y = 167.5 \text{ m}$) across the prior realizations and the reference field. The information contained in the prior realizations is visualized as a combination of a most expected mean field (blue line) and a 90 % confidence interval (blue area). The width of the 90 % confidence interval is nearly uniform across both section lines (i.e., the influence of nearby sampling points is not immediately visible). The confidence interval of the north–south and east–west section encloses the reference field at, respectively, 44 (73.33 %) and 53 (88.33 %) blocks (Fig. 6). Considering the entire grid, 3225 out of 3600 blocks (89.58 %) are characterized by a confidence interval enclosing the reference value.

Figure 7a, b displays, respectively, a north–south ($x = 57.5 \text{ m}$) and east–west section ($y = 167.5 \text{ m}$) across the AE (purple solid) and SD (orange dashed) field. A comparison between Figs. 6 and 7 shows that a severe misalignment between the reference field and confidence interval directly results in an AE which is consistently larger than

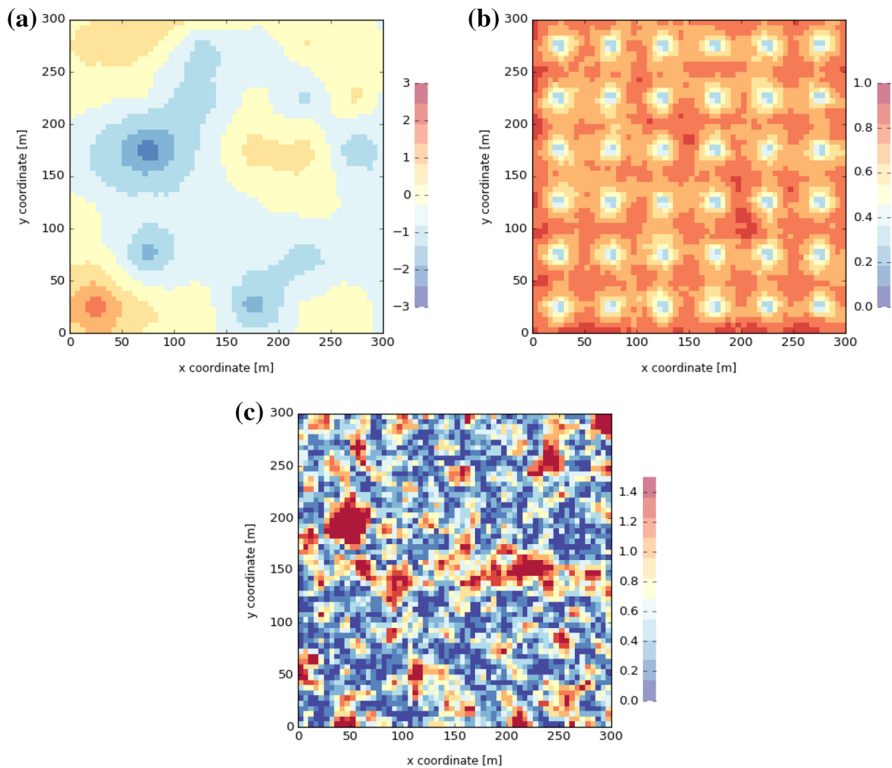


Fig. 5 Grid with 60 x 60 blocks of size 5 m × 5 m color coded according to the computed statistic **a** mean, **b** standard deviation, **c** absolute error

1.0. The three peaks of large AE in the north–south section (Fig. 7a) result from a severe, but local estimation bias (overestimation around $y = 150$ m and $y > 250$ m, underestimation at $175 \text{ m} < y < 225$ m, Fig. 6a). The two smaller peaks of AE in the east–west section (Fig. 7b) originate from a relatively large region with a smaller estimation bias (consistent underestimation between $125 \text{ m} < x < 250$ m). The SD, which provides an alternative representation of remaining uncertainty, does not change significantly across both selected sections (slight decrease in the proximity of sampling points, Fig. 7b).

4.4 Experimental scenario

A mining operation with two extraction points of unequal production rate is simulated and monitored. Figure 8 displays the outline of extraction zone I ($40 \text{ m} \times 120 \text{ m}$) and zone II ($120 \text{ m} \times 40 \text{ m}$). Each extraction zone has an area of 4800 m^2 and contains 192 blocks. From this point forward, it will be assumed that the previous constructed block realizations have a thickness of 3 m . Considering a rock density of 2.8 t/m^3 , it is not unreasonable to assume that each 75 m^3 block is transported by a truck with a payload

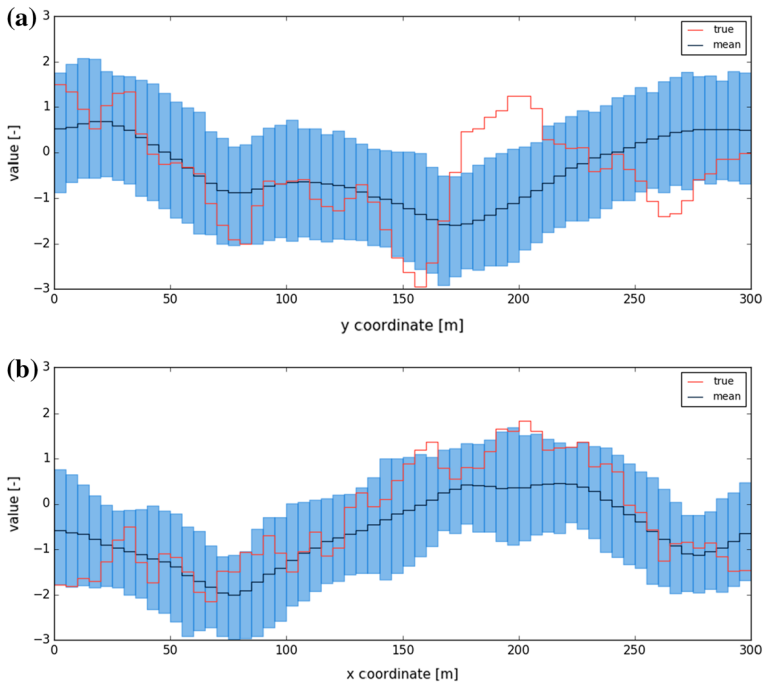


Fig. 6 Sections across prior realizations and reference field **a** north–south section at $x = 57.5$ m, **b** east–west section at $y = 167.5$ m

capacity of 210 t . Hence, the terms block and truckload are used interchangeably. The duration of the experiment equals the time T needed to excavate and process all 384 blocks from both extraction zones. This total time is subdivided into 12 discrete timesteps dt of equal length ($T = 12 * dt$). Since the total production rate is kept constant, 32 truckloads are processed during a single timestep.

Figure 9 provides a schematic overview of the extraction sequence and monitoring setup (forward simulator). The extraction sequence is visually prescribed through the numbers inside the demarcated digging blocks. During the first half time step ($[0, dt/2]$), a small digging block is extracted from mining zone I (1a, $[50 \text{ m}, 150 \text{ m}] \times [40 \text{ m}, 160 \text{ m}]$, 4 truckloads). Simultaneously, a large digging block is removed from mining zone II (1a, $[150 \text{ m}, 130 \text{ m}] \times [165 \text{ m} \times 170 \text{ m}]$, 12 truckloads). During the second half timestep ($[dt/2, dt]$), a small (1b, $[40 \text{ m}, 150 \text{ m}] \times [50 \text{ m}, 160 \text{ m}]$) and large (1b, $[165 \text{ m}, 130 \text{ m}] \times [180 \text{ m}, 150 \text{ m}]$) digging block are, respectively, extracted from mining zone I and II. A similar generic approach applies for the remaining timesteps. Generally, during each half timestep, 4 and 12 truckloads are extracted always from complementary mining zones (4 from one, 12 from the other). As a result, the total production rate remains constant: 16 truckloads per half timestep. Figure 9, however, illustrates that the local production rates (4 or 12 truckloads per half timestep) reverse after every 3^{th} timestep.

The monitoring setup is strongly related to the extraction sequence. Every timestep yields two observations, each an average of 16 blended truckloads. The composition

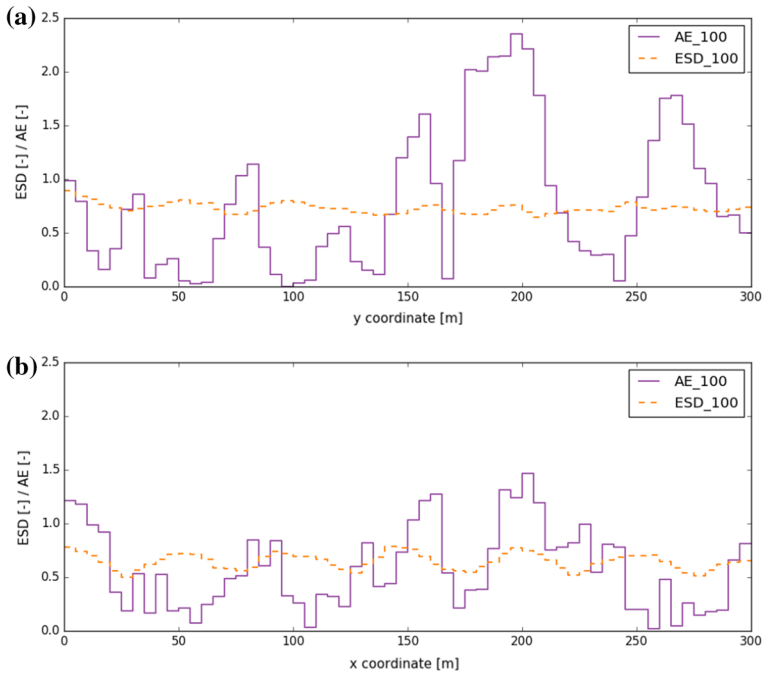


Fig. 7 Sections across absolute error (AE, purple solid) and standard deviation (SD, orange dashed) field, **a** north–south section at $x = 57.5$ m, **b** east–west section at $y = 167.5$ m

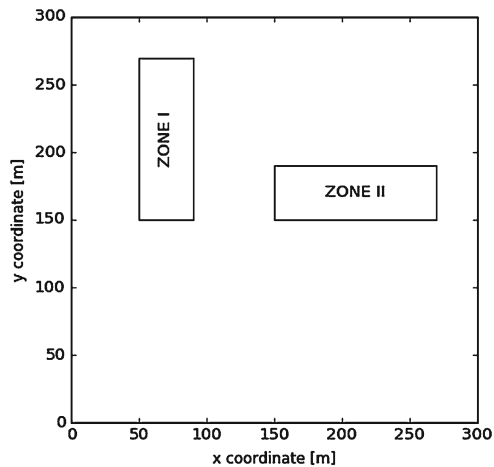


Fig. 8 Outline of extraction zone I ($[50 \text{ m}, 150 \text{ m}] \times [90 \text{ m}, 270 \text{ m}]$) and II ($[150 \text{ m}, 150 \text{ m}] \times [270 \text{ m}, 190 \text{ m}]$)

of the blend is prescribed by the defined extraction sequence. This is at time t , two observations become available, characterizing the blend resulting from the extraction during respective periods $[t - 1, t - 1 + dt/2]$ and $[t - 1 + dt/2, t]$. A forward simulator is implemented based on the described extraction sequence and measurement setup.

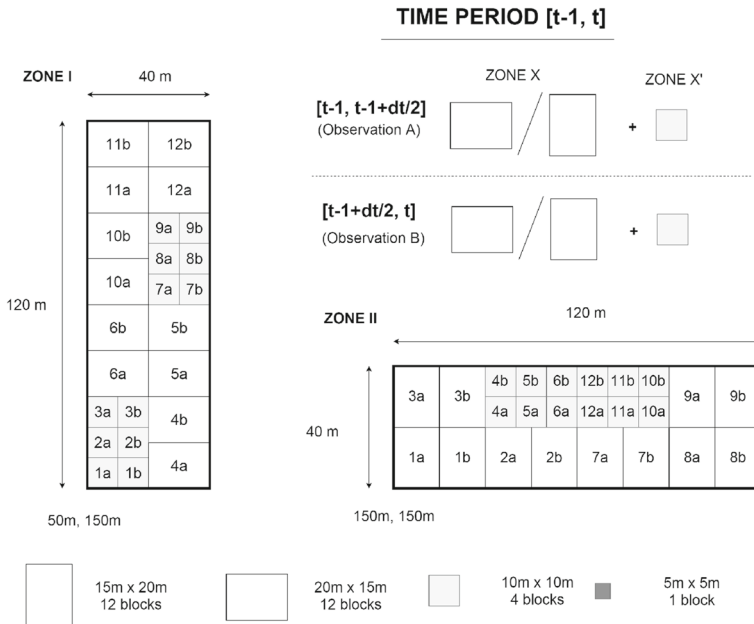


Fig. 9 Schematic overview of extraction sequence and monitoring setup: during each half timestep dt , 4 and 12 blocks are extracted from complementary zones (X and X') and blended

The forward simulator $\mathcal{A}_t$ translates each (prior) realization $\mathbf{Z}_{t-1}(:, i)$ into a vector $\mathbf{B}_t(:, i)$ with two observations (Fig. 1).

Obviously in practice, the real observation vector $\mathbf{d}_t$ results from actual measurements with inaccurate sensors. However, due to the synthetic nature of the experiment, a second forward simulator $\mathcal{A}_t^*$ needs to be built to mimic the behavior of a real monitoring network. This simulator generates inaccurate but true observations $\mathbf{d}_t$ based on the reference field $\mathbf{Z}^*$ ($\mathbf{d}_t = \mathcal{A}_t^*(\mathbf{Z}^*)$). The previous constructed simulator $\mathcal{A}_t$ is adjusted to suit this purpose. Two main modifications are made. First, the reference field $\mathbf{Z}^*$ is propagated through the simulator and not one of the realizations $\mathbf{Z}_{t-1}(:, i)$. Second, random noise is added onto the resulting observations to mimic the behavior of an inaccurate sensor. The random noise is drawn from a normal distribution with zero mean and a standard deviation of 0.0625 ($0.25/\sqrt{16}$). The reported standard deviation defines the measurement accuracy over a volume of 16 truckloads (1200 m^3). The applied sensor has a rather low accuracy. The corresponding accuracy of one truckload measurement would be 0.25, merely a quarter of the accuracy of a block estimate (block standard deviation varies between 0.283 and 0.907). In real applications, the observation and related accuracy would result from a time average of a continuous sensor response.

The described experimental scenario is executed in 29 s on a standard laptop (Intel Core i7-3520M CPU @ 2.90Ghz, 8Gb RAM memory). A single update is completed in less than 2.5 s. The computation time is expected to only slightly increase when applying this technique to ongoing operations. It is important to understand that the

time needed to solve the updating equations is only dependent on the number of grid nodes in the neighborhood subset. The size of the neighborhood subset is more determined by the geological scale of correlation and the number of extraction locations than by, e.g., the total number of nodes in a typical application problem (can be up to a billion). Larger grids though will impact the time needed to construct the neighborhood subset. Considering L extraction locations and a space-partitioning search tree, the time to construct a neighborhood subset will be in the order of $L \log_2(N)$. Therefore, it is to be expected that the time needed to complete a single update in real applications will be an order of magnitude smaller than the time needed to collect time-averaged measurements (a few minutes versus a few hours).

5 Results

The performance of the algorithm is evaluated using three sets of criteria. (1) Fields and cross sections are visually inspected. Mean and reference fields are compared, the reliability of the 90 % confidence interval is assessed and the evolution of the AE and SD is studied. (2) Global assessment statistics are computed. The overall quality of a set of realizations is evaluated by the magnitude of two single value measures: the root mean square error (RMSE) and the spread. (3) Local assessment statistics are computed for 14 different areas. RMSE and spread are calculated per local area to aid a quantitative comparison. Scatter plots, box plots and empirical probability plots are drawn to visually evaluate local improvements.

5.1 Visual inspection

Figure 10 displays how the mean field changes through time when assimilating sensor observations from the production chain. The two large rectangular outlines demarcate the extent of both extraction zones (I and II). The four smaller rectangles delineate the digging block which were extracted, blended and measured during the considered time interval. A comparison of Fig. 10a–l with the prior mean field (Fig. 5a) and the reference field (3a) indicates that the mean field (the current best estimate) continuously improves in and around both extraction zones.

Figure 11a, b depict, respectively, a north–south ($x = 67.5$ m) and east–west ($y = 167.5$ m) section across the reference field and final realizations (after 12 updates). Both figures support the previous conclusion. The mean field in the neighborhood of both extraction zones ($150 \text{ m} < y < 270 \text{ m}$ and $150 \text{ m} < x < 270 \text{ m}$) considerably improved compared to the initial model (Fig. 6). Two out of the three occurrences of severe local estimation bias in the north–south section are corrected ($175 \text{ m} < y < 225 \text{ m}$ and $y > 250 \text{ m}$, Fig. 6a and 11a). The third occurrence, an overestimation bias around $y = 150 \text{ m}$, could not be corrected. This local deviation did go undetected because large positive deviations of small digging blocks were blended with smaller negative deviations of larger digging blocks (1a and 1b from zone I mixed with 1a and 1b from zone II, Fig 9). The two local corrections along the north–south section can also be observed in Fig. 10f, j (increase mean field, $175 \text{ m} < y < 225 \text{ m}$) and 10k (decrease mean field, $y > 250 \text{ m}$). The algorithm further corrected the underestimation bias in

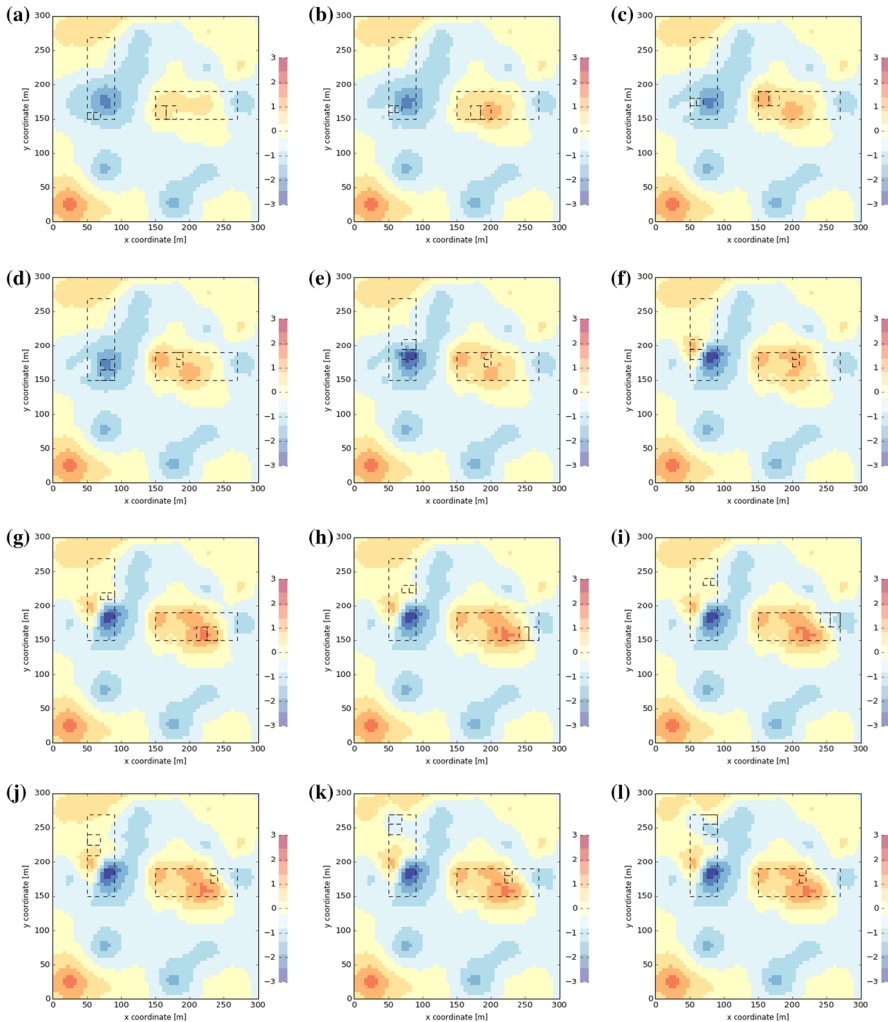


Fig. 10 Evolution of mean field through time, from update 1 (a) to 12 (l)—grid with 60×60 blocks of size $5 \text{ m} \times 5 \text{ m}$ color coded according to the computed average block value

the east–west section (Figs. 6b – 11b, $125 \text{ m} < x < 250 \text{ m}$). The observed correction in the east–west section corresponds to a higher mean field in mining zone II (e.g., compare Fig. 10a, l).

Figure 12a, b displays a north–south ($x = 57.5 \text{ m}$) and an east–west ($y = 167.5 \text{ m}$) section across the initial and final AE field. The 12 updates result in a significant reduction of the absolute error around the extraction zones. Two of the three peaks of large AE in the north–south section have disappeared (Fig. 12a, $175 \text{ m} < y < 225 \text{ m}$ and $y > 250 \text{ m}$). In the corresponding areas, a reduction in AE from 2.0 to 0.5 is observed. The height of the third peak, however, slightly increases from about 1.6 to 2.3 (Fig. 12a $y = 150 \text{ m}$). A plausible explanation for this behavior has been

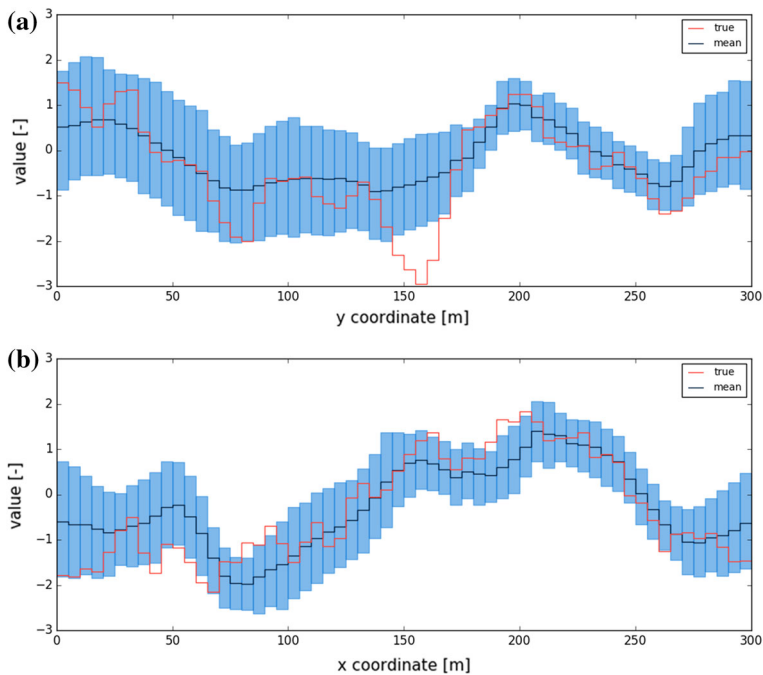


Fig. 11 Sections across updated realizations (after update 12) and reference field: **a** north–south section at $x = 57.5$ m, **b** east–west section at $y = 167.5$ m

provided previously. The correction of the underestimation bias around mining zone II causes a reduction of AE in the eastern part of the east–west section (Fig. 12b, $125 \text{ m} < x < 250 \text{ m}$). The increase of AE in the western part of this section is directly related to the discussed local underperformance of the algorithm at the southern border of mining zone I (Fig. 10a $y = 167.5$ m, Fig. 10b $x = 57.5$ m).

The assimilation of sensor observations results as well in an adjustment of the uncertainty model. The width of the 90 % confidence interval is approximately halved within both extraction zones (compare Figs. 6 and 11, $150 \text{ m} < y < 270 \text{ m}$ and $< 150 \text{ m} < x < 270 \text{ m}$). The remaining uncertainty inside the excavated areas is attributed to a combination of measurement error and scale of support effects. In other words, average inaccurate measurements of blended truckloads are not sufficient to uniquely characterize individual blocks without any remaining uncertainty. The evolution of the SD field through time illustrates that the remaining uncertainty is not uniformly adjusted across the excavated areas (Fig. 13). The updating algorithm seems to assign a lower uncertainty (lower SD) to the larger digging blocks. The SD of updated blocks in large digging areas varies between 0.25 and 0.45 (Fig. 13). The SD of updated blocks in small digging areas is significantly higher, and values between 0.55 and 0.75 are observed. A comparison of the SD fields across Fig. 13 indicates that the adjustment of the uncertainty model is not only limited to already excavated areas. For example, Fig. 13a, b shows that the uncertainty near the middle of mining zone II reduces from somewhere around 0.7 to 0.6 over even 0.5. This

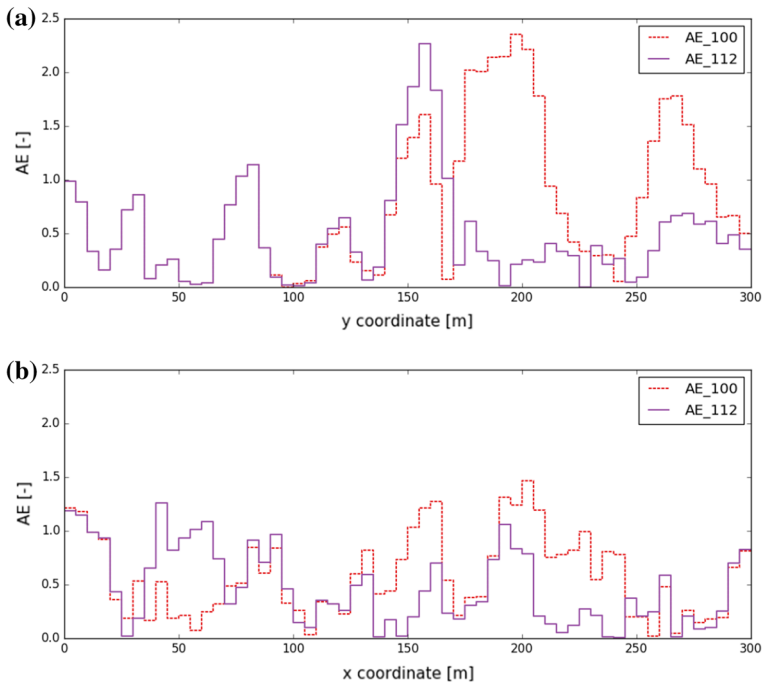


Fig. 12 Sections across absolute error before (*dashed red*) and after 12 updates (*solid purple*). **a** North–south section at $x = 57.5$ m, **b** east–west section at $y = 167.5$ m

reduction occurs during the second update when none of the material near the middle of mining zone II has already been excavated. Figs. 13 and 5b illustrate that a reduction in uncertainty is observed up to a distance of 30 m from the border of a digging block.

Figure 14a, b depicts, respectively, a north–south ($x = 67.5$ m) and east–west ($y = 167.5$ m) section across the prior (Fig. 5b) and final (13l) SD field. The north–south section shows that the reduction of the SD is lower in the smaller digging areas ($150 \text{ m} < y < 180 \text{ m}$) compared to the larger ones ($180 \text{ m} < y < 270 \text{ m}$). North of the northern border of extraction zone I ($y > 270 \text{ m}$), the SD increases from 0.25 (inside, $180 \text{ m} < y < 270 \text{ m}$) towards 0.75 over a distance of about 30 m. The uncertainty as well reduces up to 30 m away from the southern border ($120 \text{ m} < y < 150 \text{ m}$), though this reduction is less remarkable due to the larger SD values in the smaller digging areas ($150 \text{ m} < y < 180 \text{ m}$). The east–west section primarily intersects large digging areas (except for $50 \text{ m} < x < 70 \text{ m}$); hence, the SD inside both extraction zones generally reduces to values between 0.25 and 0.35 (Fig. 14a). The SD inside the smaller digging areas ($50 \text{ m} < x < 70 \text{ m}$) reaches values near 0.65. The effects of a reduction in uncertainty are again observed up to 30 m away from the boundaries of both extraction zones (zone I, $x < 50 \text{ m}$ and $x > 90 \text{ m}$; zone II, $x < 150 \text{ m}$ and $x > 270 \text{ m}$). The differences between the prior and final SD field displayed in the east–west section (Fig. 14b) are less striking compared

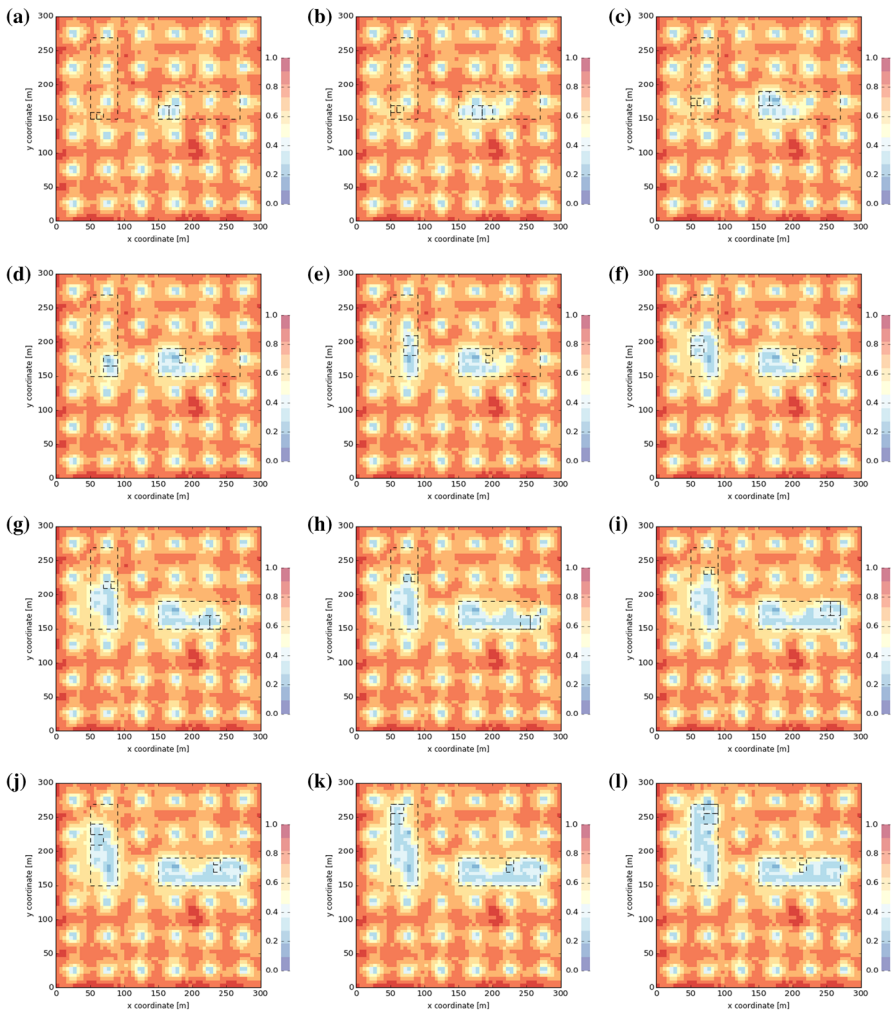


Fig. 13 Evolution of SD field through time, from update 1 (a) to 12 (l)-grid with 60×60 blocks of size $5 \text{ m} \times 5 \text{ m}$ color coded according to the computed block standard deviation

to those in the north–south section (Fig. 14a). This is due to lower and more variable prior SD field in the east–west section (larger effect of nearby sampling points) and has nothing to do with the updating results which are similar for both cross sections.

5.2 Global assessment statistics

The RMSE and spread are single value measures expressing the global quality of a set of realizations. As such, they allow for a more convenient comparison of realization sets across different timesteps. The RMSE measures the overall deviation between the

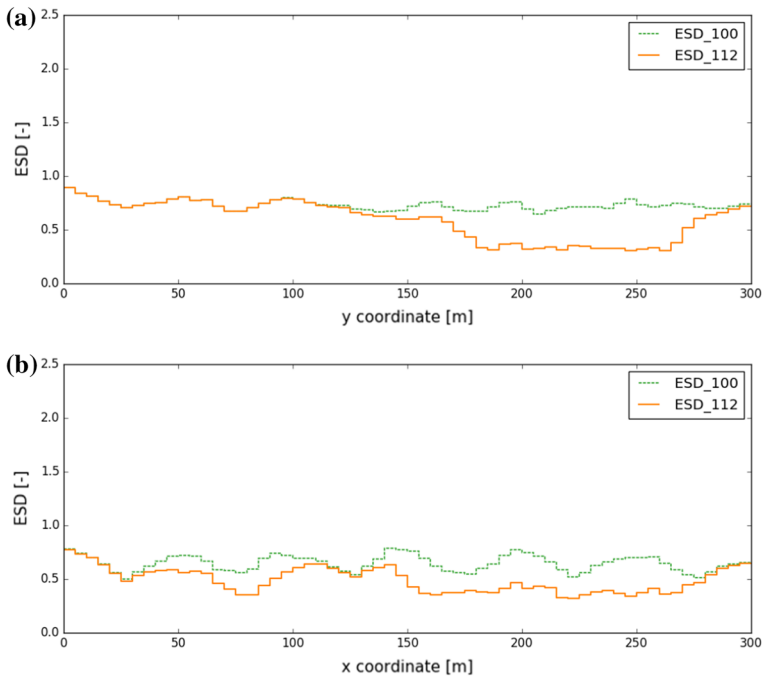


Fig. 14 Sections across standard deviation before (*dashed green*) and after 12 updates (*solid orange*). **a** North–south section at $x = 57.5$ m, **b** East–West section at $y = 167.5$ m

reference field ($\mathbf{Z}^*$) and the mean field ($\hat{\mathbf{Z}}_t$) at a given timestep. The RMSE is defined as the root of the average squared difference between all available true and estimated block values

$$\text{RMSE}_t = \sqrt{\frac{1}{N} \sum_{n=1}^N [\mathbf{Z}^*(n) - \hat{\mathbf{Z}}_t(n)]^2}. \quad (20)$$

The spread summarizes the overall estimated uncertainty and is calculated as the root of the average block variance

$$\text{SPREAD}_t = \sqrt{\frac{1}{N} \sum_{n=1}^N \text{VAR}(\mathbf{Z}_t(n, :))}. \quad (21)$$

In both formulations (Eqs. 20 and 21), n iterates over all 3600 blocks in the grid. Ideally, the RMSE reduces in function of the number of incorporated measurements. During the course of the experiment, the uncertainty model is adjusted as well to account for the remaining error. Hence, the need for a spread which approximates the RMSE.

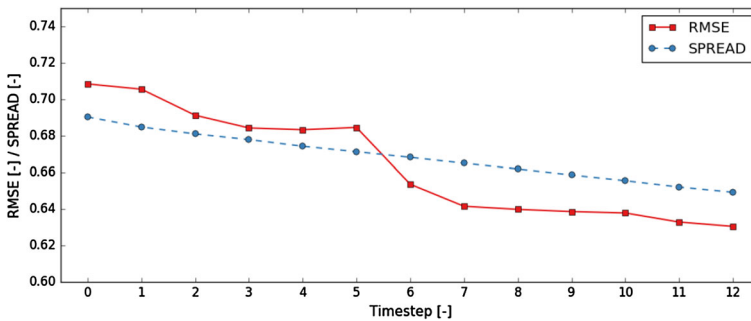


Fig. 15 Evolution of RMSE and spread through time

Figure 15 shows the evolution of the global RMSE and spread through time. The RMSE of the prior realizations amounts to 0.708, while the prior spread is 0.690. Over the course of the experiment, the global RMSE drops by about 11 % (from 0.708 to 0.630). The spread decreases accordingly, a reduction of 5% is observed (from 0.690 to 0.649). Figure 15 indicates that the spread decreases at a constant rate. This was to be expected since a constant material volume is measured during each timestep. There is no reason to assume that the information value of similar observations would be different over time. Figure 15 illustrates that the RMSE tends to reduce at the same rate. The more volatile behavior is possibly caused by its large sensitivity to local improvements. For example, the rate of reduction in RMSE from timestep 5 to 6 is exceptionally large and inconsistent with the general pattern. This exceptionally large reduction originates from a local correction of an anomalous area. Figure 5c already indicated that the prior AE was consistently larger than 2.0 in an area bounded by [25 m, 175 m] $\times$ [75 m, 225 m]. During the 6th timestep, a significant part of this area is excavated and measured. The algorithm detects this local anomaly and performs a correction. Figure 10e, f clearly shows that the mean field in this local area is lifted during the 6th timestep.

5.3 Local assessment statistics

The previously reported global statistics are computed over all available 3600 blocks. During an update, only a fraction of this number is adjusted. The effects of local improvements are thus smeared out when computing a measure of global quality. The following presents a more local assessment to better comprehend the effects of an updating sequence.

A total of 14 local areas are defined and displayed in Fig. 16. Two local areas (I0 and II0) coincide with the previous demarcated extraction zones (zone I and II, Fig. 8). Six constitute frames of 5m thickness around mining zone I (I1 to I6). Another six are laid out around mining zone II (II1 to II6).

For each local area, a true value is computed by averaging blocks of the reference field located inside the corresponding boundary. Similarly, each realization is converted into 14 local area values. Upon completion, 200 simulated local area values are

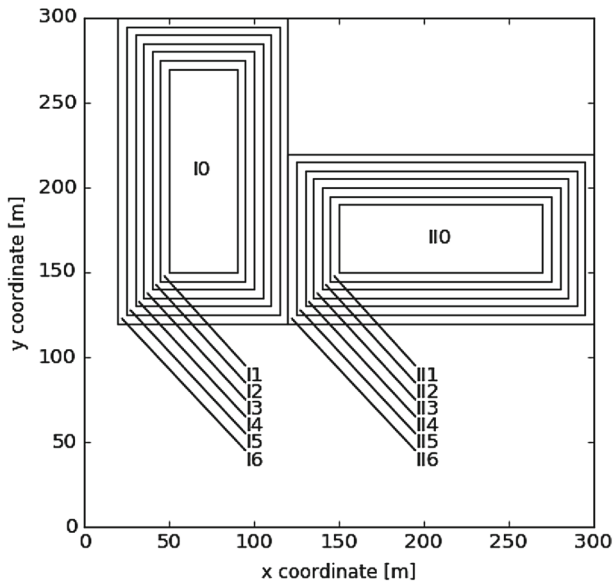


Fig. 16 Outline of 14 local areas, 2 extraction zone of 4800 m^2 , and 12 surrounding frames with a thickness of 5 m

collected and further used to compute a best estimate (i.e., mean) and quantile values. Figure 17 displays the true local area values versus their best estimates before and after 12 updates. A zero distance between a point and the bisector refers to a perfect alignment of a true and estimated value. The scatter plot indicates that the quality of the prior model (blue dots) is location dependent. In general, local area estimates deviate more from their reference values in and around mining zone II (open markers) than they do in and around mining zone I (filled markers). The 12 updates alter the realizations such that the estimated area values in and around zone II are shifted towards the bisector (red squares, Fig. 17). This corresponds to a general increase in the mean field in the vicinity of zone II (compare Fig. 10a, I). Notice the large change in area value II0 and II1, both change respectively from 0.062 to 0.547 (0.673) and from -0.199 to 0.236 (0.427). Their corresponding reference values are indicated between brackets. The best estimates in mining zone I do not experience such drastic changes.

Figure 18 depicts box plots constructed from the computed local area quantiles. The box plots indicate that the prior uncertainty model partially accounts for the previously observed initial deviations (Fig. 18a). For about 10 of the 14 local areas, the reference field is located inside the 90 % confidence interval. The mismatch between the reference value and confidence interval is remarkable in areas II0 to II4. After updating 12 times, the box plots characterizing the local areas in mining zone II are shifted upwards and become shorter. Those characterizing mining I undergo a similar transition, although their upward shift is much more subtle. A uniform shift refers to a correction of bias, while shorter box plots indicate a reduction of uncertainty.

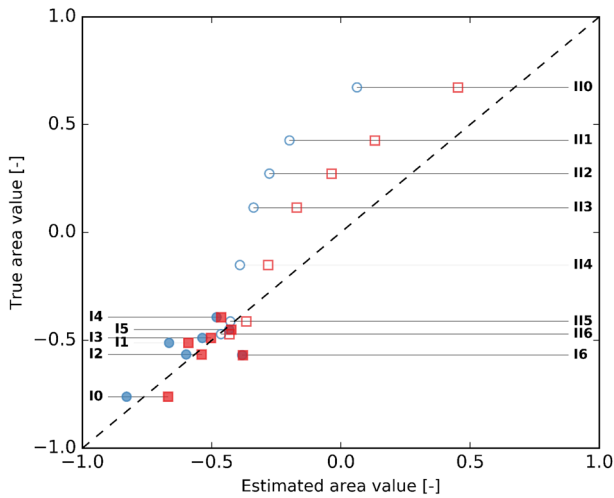


Fig. 17 Scatterplot of true versus estimated area values before (*blue dots*) and after 12 updates (*red squares*)—filled markers refer to areas in mining zone I, open markers refer to areas in mining zone II

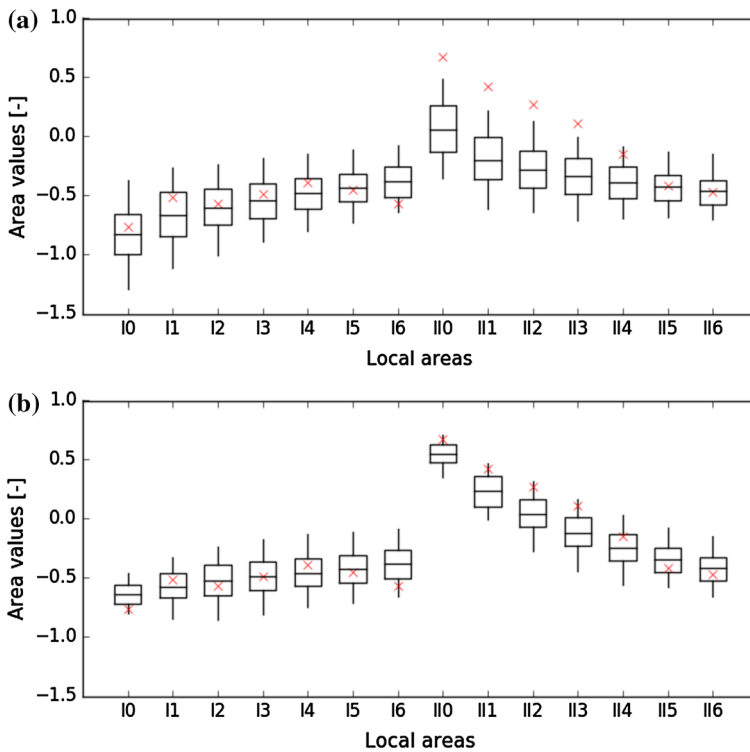


Fig. 18 Box plots constructed from the computed local area quantiles at timestep 0 (a) and timestep 12 (b)—reference values are indicated as *red crosses*, *box* and *whiskers* refer, respectively, to the 50 and 90 % confidence interval

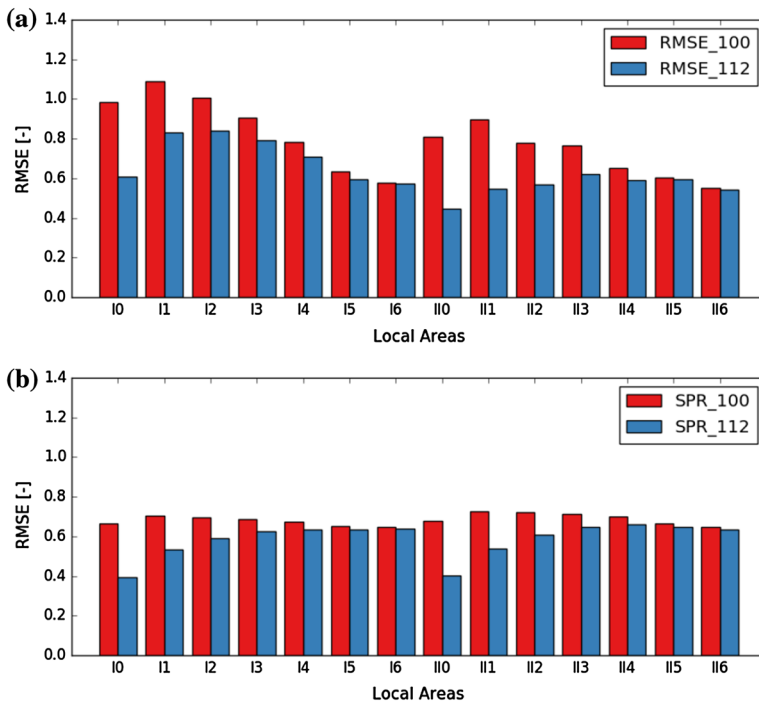


Fig. 19 Bar charts displaying the local area RMSE (a) and spread (b) of the prior (red) and updated realizations (blue)

The previously defined RMSE and spread can also be computed locally. In Eqs. 20 and 21, n then refers to blocks confined to one of the specified local areas. Figure 19 displays the prior and posterior RMSE and spread per local area. Apparently, the prior RMSE of the local areas in mining zone I is slightly larger than those in mining zone II (Fig. 19a). This might seem contradictory, since previously larger deviations between local area estimates and reference values were observed in mining zone II, and not in zone I (Fig. 17). The contradiction directly results from the effect of different bias mechanisms on the computed statistics. In mining zone I, three occurrences of a severe local under- or overestimation of the prior model are observed (zone I, Figs. 3a, 5a and 10l). When computing an area estimate, the local deviations are averaged out. However due to the squared difference in Eq. 18, these local large deviations do have a significant impact on the resulting RMSE value. The prior model in zone II is affected by a medium single bias extending over a larger area (zone II, Figs. 3a, 5a and 10l). This regional bias is not averaged out, and hence the larger deviation between an area estimate and the reference value. Since the regional bias is of medium magnitude, the resulting RMSE is not as large as that of mining zone I.

After 12 updates, the RMSE of local areas I0 and II0 drop by about 38 and 45 % (I0: from 0.984 to 0.608, II0: from 0.811 to 0.447, Fig. 19a). The first three local areas around mining zone I (I1, I2, I3, Fig. 16) observe a reduction in error of about 24, 17 and 13 %. The posterior RMSE values are, respectively, 0.833 (I1), 0.839 (I2)

and 0.790 (I3). In the three frames surrounding mining zone II (II1, II2, II3, Fig. 16), the RMSE drops by about 39, 27 and 19 %, reaching posterior RMSE values of 0.547 (II1), 0.569 (II2) and 0.620 (II3). Around mining zone II, a significant improvement is observed and small posterior RMSE values are attained (0.55 to 0.60). The updates are still successful around the mining zone, although medium RMSE values are obtained (0.80).

Figure 19b displays the prior and posterior spread per local area. The prior spread is more or less uniform across the areas and varies between 0.647 and 0.724. After 12 updates, the spread of local areas I0 and II0 drop by about 40 and 41 % (I0: from 0.663 to 0.396, II0: from 0.677 to 0.401). The first local areas surrounding both extraction zones (I1 and II1) observe a reduction in spread of about 24 and 26 % (I1: from 0.704 to 0.533, II1: from 0.724 to 0.539). The spread in local areas I2 and II2 reduces by about 15 and 16 % (I2: from 0.696 to 0.591, II2: from 0.720 to 0.607). A reduction of 9 and 10 % is observed in local areas I3 and II3 (I3: from 0.687 to 0.624, II3: from 0.714 to 0.646). The reported spread values indicate that the uncertainty gradually increases in function of the distance from the extraction zones. The initial bias does not seem to affect the posterior spread values. Further, notice the similarities between the reduction of RMSE in mining zone I and the overall spread reduction.

To conclude the discussion of the experimental results, global and local empirical probability curves are presented and analyzed. For each block included in the analysis, empirical quantiles are computed according to a predefined set of theoretical probabilities (e.g., 0.025–0.975). For a given theoretical probability, corresponding block quantiles are retrieved and compared against the reference values. The proportion of reference values not exceeding the block quantiles is referred to as the empirical probability. The procedure is repeated for all predefined theoretical probabilities. The interested reader is referred to Olea (2012). A good uncertainty model would result in probability pairs located along the bisector. This would for example imply that for 10 % of the investigated blocks (empirical probability), the reference value is lower than the corresponding 10 % quantile (the 10 % here refers to the theoretical probability).

Figure 20a displays the global empirical probability curves computed from prior and updated realizations. The curve confirms a globally good agreement between deviations (error) and uncertainty (spread) in the prior model (due to medium dense sampling grid). Updating results in an even better alignment of the global probability curve with the bisector.

Figure 20b–d displays two prior and two posterior local probability curves computed from blocks in, respectively, local areas I0 and II0, local areas I2 and II2 and local areas I4 and II4 (Fig. 16). The figures illustrate that the regional deviations in mining zone II are not accounted for in the prior uncertainty model (purple triangles). The occurrence of low values is severely overestimated. For example in only 6 % of the blocks in area I0 is the reference value below the computed 40 % quantile value. The prior probability curves for areas II0, II2 and II4 indicate that the block distributions in and around extraction zone II are centered around values which are too low (purple triangles, Fig. 20b–d). Updating causes a significant upward shift of the corresponding distributions (orange pentagons, Fig. 20b–d). However, the posterior empirical probability curves indicate that the upward shift could have been slightly larger. Figure 18b confirms this conclusion. A slightly higher position of the box plots,

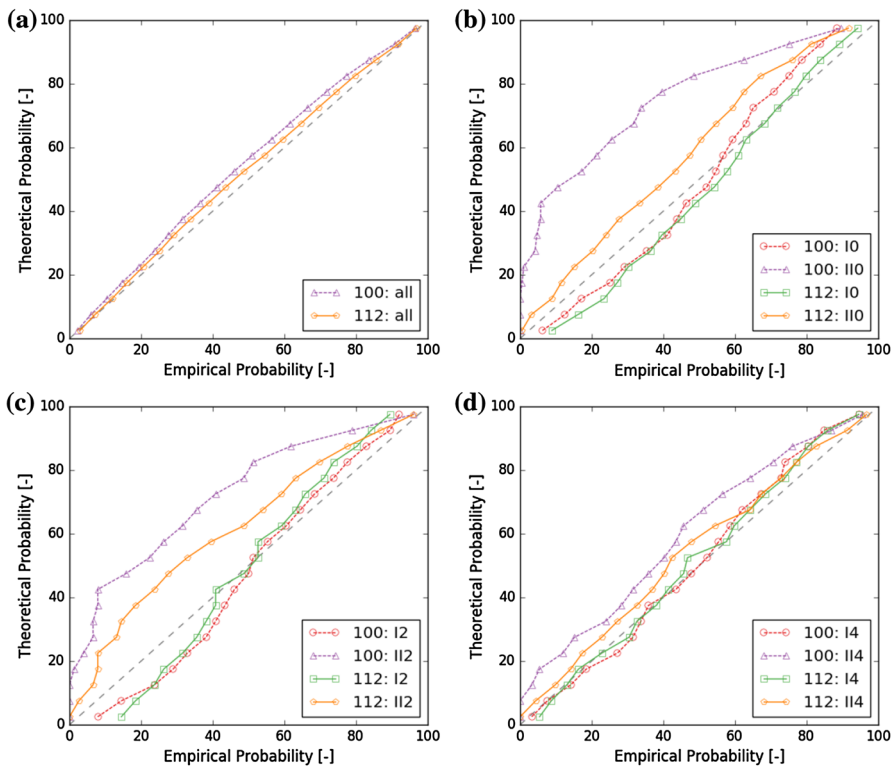


Fig. 20 Empirical versus theoretical probabilities computed over prior and updates realizations—**a** entire grid, **b** local area I0 and II0, **c** local areas I2 and II2, **d** local areas I4 and II4

representing local areas II0 to II4, would result in reference values located inside the 50% confidence interval. The prior empirical probability curves of areas I0, I2 and I4 approve the overall quality of the prior uncertainty model around mining zone I (red circles, Fig. 20b–d). A more detailed analysis reveals that the tails of the distributions are a bit too long. Such a conclusion can be drawn from the fact that the occurrence of low values and high values is slightly overestimated ('S' shape instead of straight line). Updating generally results in a better alignment of the posterior curves with the bisector (green squares, Fig. 20b–d). This corresponds to shorter confidence interval, as observed in Fig. 18b.

6 Conclusion

This contribution presents a new algorithm to automatically assimilate online data from a production monitoring network into the grade control model. The practicality of the algorithm mainly results from the design decision to exclude the forward observation model from the computer code. Instead, for each unique application, a case-specific forward simulator is built and used independently from the existing code. This results

in great flexibility, allows for a better integration of expert knowledge and facilitates interdisciplinary cooperation.

Once built, the forward simulator is run to translate grade control realizations into observation realizations. The resulting realization sets are subsequently used to compute the empirical covariances. These covariances mathematically describe the link between the sensor observations and blocks from the grade control model. There is no need to formulate and linearize an analytical forward observation model, let alone compute its inverse. The forward simulator further ensures that the distribution of the Monte Carlo samples already reflect the support of the concerned random values. As a result, the necessary covariance, derived from these Monte Carlo samples, inherently account for differences in the scale of support.

A Gaussian anamorphosis option is implemented to deal with suboptimal conditions related to non-Gaussian distributions. A specific algorithm structure ensures that the sensor precision (measurement error) can be defined on its original units and does not need to be translated into a normal score equivalent.

When due to practical considerations the Monte Carlo samples are rather limited, numerical inaccuracies may arise. An interconnected parallel updating sequence (helix) can be configured to reduce the effects of filter inbreeding. Due to the empirical computation of the covariance, degrees of freedom are lost over time. Eventually, this might result in a collapse of covariances. A neighborhood option is implemented to constrain computation time and memory requirements. Since observations are collected on blended material streams originating from multiple extraction zone, different neighborhoods need to be considered simultaneously. A neighborhood strategy also partially reduces the effect of spurious correlations. Two covariance correction options are implemented to contain the propagation of statistical sampling errors originating from the empirical computation of covariances. The localization correction is based on the assumption that the accuracy of a covariance estimate decreases with distance from an extraction point. Hence, grid anomalies are tapered around central extraction locations. The error correction uses an off-line Monte Carlo simulation to precompute the accuracy of a certain correlation estimate. Corresponding correction factors are stored in lookup table and used to adjust the elements of the forecast error covariance matrix.

An artificial experiment is conducted to showcase the capabilities of the developed algorithm. A mining operation with two extraction points of unequal production rate is simulated. The resulting material streams are blended and inaccurate observations are made. A total of 12 updates is performed, each one based on 2 inaccurate observations of 16 blended blocks. The results indicate a 38–45 % improvement in RMSE inside both extraction zone. Improvements ranging from 39 to 13 % were observed in a close neighborhood. The results are very promising, especially considering the large measurement volume and low sensor accuracy.

Future research should focus on the influence of the monitoring network and geological environment on the performance of the algorithm. A series of experiments needs to be conducted to measure the effects of different sensor precisions, measurement volumes, update intervals and blending ratios. The results of such a study could be used to formulate recommendations for designing an optimal monitoring network, guaranteeing optimal algorithm performance. The authors expect that for

a given blending ratio the order of extraction does not significantly impact the final results. This assumption has not been proven and should be verified.

It is also important to understand the limitations of the algorithm regarding the geological environment and the prior model hereof. The question still remains whether good performance can be achieved when the prior model and the reality differ substantially. Also, the effects of different spatial correlation structures are not yet understood (different covariance models, ranges, nugget,...). Another possible point of attention could be the further improvement of the localization error correction option. The current implementation (Gaspari–Cohn correlation function) results in a differential covariance correction of blocks inside the same digging area. Blocks near the extraction point receive a lower correction factor than blocks slightly further away. Especially when considering large measurement volumes (digging blocks), a tapering function with a plateau might be more suitable. The function should be defined such that all blocks inside the digging area are not corrected (correction factor of 1) and that tapering starts from the border of the digging areas.

Open Access This article is distributed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license, and indicate if changes were made.

References

- Alabert F (1987) The practice of fast conditional simulations via the lu decomposition of the covariance matrix. *Math Geol* 19:368–386
- Anderson J (2012) Localization and sampling error correction in ensemble kalman filter data assimilation. *Am Meteorol Soc* 140:2359–2371
- Beal D, Brasseur P, Brankart JM, Ourmieres Y, Verron J (2010) Characterization of mixing errors in a coupled physical biogeochemical model of the north atlantic: implications for nonlinear estimation using gaussian anamorphosis. *Ocean Sci* 6:247–262
- Benndorf J (2013) Application of efficient methods of conditional simulation for optimizing coal blending strategies in large continuous open pit mining operations. *Internat J Coal Geol* 122:141–153
- Benndorf J (2015) Making use of online production data: sequential updating of mineral resource models. *Math Geosci* 47:547–563
- Bertino L, Evensen G, Wackernagel H (2002) Combining geostatistics and kalman filtering for data assimilation in an estuarine system. *Inverse Probl* 18:1–23
- Burgers G, van Leeuwen P (1998) Analysis scheme in the ensemble Kalman filter. *Mon Weather Rev* 126:1719–1724
- Chen Y, Oliver D (2010) Cross-covariances and localization for enkf in multiphase flow data assimilation. *Comput Geosci* 14:579–601
- Cheng C, Parzen E (1997) Unified estimators of smooth quantile and quantile density functions. *J Stat Plan Inference* 59:291–307
- Chiles JP, Delfiner P (2012) *Geostatistics modeling spatial uncertainty*. John Wiley and Sons, Hoboken, New Jersey
- Cohn S, da Silva A, Guo J, Sienkiewicz M, Lamich D (1998) Assessing the effects of data selection with the dao physical space statistical analysis system. *Mon Weather Rev* 126:2913–2926
- Dagan G (1985) Stochastic modeling of groundwater flow by unconditional and conditional probabilities, 2, the inverse problem. *Water Resour Res* 21:65–72
- Davis M (1987) Production of conditional simulations via the lu triangular decomposition of the covariance matrix. *Math Geol* 19:91–98

- Deutsch C, Journel A (1992) GSLIB geostatistical software library and user's guide. Oxford University Press, Oxford
- Dimitrakopoulos R (1998) Conditional simulation algorithms for modelling orebody uncertainty in open pit optimisation. *Int J Surf Mining Reclam Environ* 112:173–179
- Dimitrakopoulos R, Godoy M (2014) Grade control based on economic ore/waste classification functions and stochastic simulations; examples, comparisons and applications. *Trans Inst Min Metall Sect A Min Technol* 123:90–106
- Dimitrakopoulos R, Jewbali A (2013) Joint stochastic optimization of short- and long-term min production planning: method and application in a large operating gold mine. *Trans Inst Min Metall Sect A Min Technol* 122:110–123
- Dowd P (1994) Risk assessment in reserve estimation and open-pit planning. *Trans Inst Min Metall Sect A Min Technol* 103:148–154
- Evensen G (1992) Using the extended kalman filter with a multi-layer quasi-geostrophic ocean model. *J Geophys Res* 97(C11):905–924
- Evensen G (1994) Sequential data assimilation with a nonlinear quasi-geostrophic model using monte carlo methods to forecast error statistics. *J Geophys Res* 99(C5):10143–10162
- Evensen G (1997) Advanced data assimilation for strongly nonlinear dynamics. *Mon Weather Rev* 125:1342–1354
- Evensen G, van Leeuwen P (1996) Assimilation of geosat altimeter data for the agulhas current using the ensemble kalman filter with a quasigeostrophic model. *Mon Weather Rev* 124:85–96
- Franssen HH, Kaiser H, Kuhlmann U, Bauser G, Stauffer F, abd W. Knizelbach, R.M., (2011) Operational real-time modeling with ensemble kalman filter of variably saturated subsurface flow including stream-aquifer interaction and parameter updating. *Water Resources Research* 47:1–20
- Gaspari G, Cohn S (1999) Construction of correlation functions in two and three dimensions. *Quater J R Meteorol Soc* 125:723–757
- Gomez-Hernandez, J., Cassiraga, E., 2000. Sequential conditional simulations with linear constraints, in: Geostatistics 2000, Cape Town, Monestiez P., Allard., D. and Froideveaux, R. (eds), Geostatistical association of Southern Africa
- Goovaerts P (1997) Geostatistics for natural resource evaluation. Oxford University Press, New York
- Hansen T, Journel A, Tarantola A, Mosegaard K (2006) Linear inverse gaussian theory and geostatistics. *Geophysics* 71:R101–R111
- Hansen T, Mosegaard K (2008) Visim: sequential simulation for linear inverse problem. *Comput Geosci* 34:53–76
- Harter T, Yeh R (1996) Conditional stochastic analysis of solute transport in heterogeneous, variably saturated soils. *Water Resour Res* 32:1597–1609
- Harvey CF, Gorelick S (1995) Mapping hydraulic conductivity: sequential conditioning with measurements of solute arrival time, hydraulic conductivity, and local conductivity. *Water Resour Res* 31:1615–1626
- Heidari, L., Gervais, V., Ravalec, M.L., Wackernagel, H., 2011. History matching of reservoir models by ensemble kalman filtering: the state of the art and a sensitivity study, in: Uncertainty analysis and reservoir modeling, AAPG memoir 96, American Association of Petroleum Geologists. pp 249–264
- Hoeksema R, Kitanidis P (1984) An application of the geostatistical approach to the inverse problem in two-dimensional groundwater modeling. *Water Resour Res* 20:1003–1020
- Houtekamer P, Mitchell H (1998) Data assimilation using an ensemble kalman filter technique. *Mon Weather Rev* 126:796–811
- Hu L, Zhao Y, Liu Y, Scheepens C, Bouchard A (2012) Updating multipoint simulations using the ensemble kalman filter. *Comput Geosci* 51:7–15
- Jansen J, Douma S, Brouwer D, van den Hof P, abd AW, Hemink OB (2012) Closed-loop reservoir management. In: Proceedings SPE reservoir simulation symposium, Woodlands, Society of petroleum engineers
- Jewbali A, Dimitrakopoulos R (2011) Implementation of conditional simulation of successive residuals. *Comput Geosci* 37:129–142
- Journel A, Huijbregts C (1978) Mining geostatistics. Academic press, London
- Kalman R (1960) A new approach to linear filtering and prediction problems. *Trans ASME J Basic Eng* 82:35–45
- Kitanidis P, Vomvoris E (1983) A geostatistical approach to the inverse problem in groundwater modelling (steady state) and one-dimensional simulations. *Water Resour Res* 19:677–690

- Lessard J, de Bakker J, McHugh L (2014) Development of ore sorting and its impact on mineral processing economics. *J Miner Eng* 65:88–97
- McCarthy DWP (1999) Start-up performance of new base metal projects. In: Adding value to the Carpentaria mineral province, Mt Isa, Qld, Australian Journal of Mining
- McCarthy P (2003) Managing technical risk for mine feasibility studies. In: Mining Risk Management conference, Sydney, The Australian Institute for mining and metallurgy
- Morley C (2014) Monitoring and exploiting the reserve. In: Mineral resource and reserve estimation - the AusIMM guide to good practice, monograph 30, 2nd edition, The Australasian Institute of Mining and Metallurgy. pp 647–657
- Nienhaus K, Pretz T, Wortuba H (2014) Sensor technologies: impulses for the raw materials industry. RWTH Aachen, Aachen
- Olea R (2012) Building on crossvalidation for increasing the quality of geostatistical modeling. *Stoch Environ Res Risk Assess* 26:78–82
- Oliver D, Reynolds A, Liu N (2008) Inverse theory for petroleum reservoir characterization and history matching. Cambridge University Press, Cambridge
- Peattie R, Dimitrakopoulos R (2013) Forecasting recoverable ore reserves and their uncertainty at morilla gold deposit, mali: an efficient simulation approach and future grade control drilling. *Math Geosci* 45:1005–1020
- Rendu JM (2002) Geostatistical simulations for risk assessment and decision making: the mining industry perspective. *Int J Surf Min Reclam Environ* 16:122–133
- Rubin Y, Dagan G (1987) Stochastic identification of transmissivity and effective recharge in steady groundwater flow, I, theory. *Water Resour Res* 23:1185–1192
- Sakov P, Betino L (2011) Relations between two common localisation methods for the enfk. *Comput Geosci* 15:225–237
- Schoniger A, Nowak W, Franssen HJH (2012) Parameter estimation by ensemble kalman filters with transformed data: approach and application to hydraulic tomography. *Water Resources Research* 48:4502–4420
- Schewpe F (1973) Uncertain dynamic systems: modelling, estimation, hypothesis testing, identification and control. Prentice-Hall, Englewood Cliffs
- Simon E, Bertino L (2009) Application of the gaussian anamorphosis to assimilation in a 3-d coupled physical-ecosystem model of the north atlantic with the enfk: a twin experiment. *Ocean Sci* 5:495–510
- Sun N, Yeh WG (1992) A stochastic inverse solution for transient groundwater flow: parameter identification and groundwater flow: parametric identification and reliability analysis. *Water Resour Res* 28:3269–3280
- Tarantola A (2005) Inverse problem theory and methods for model parameter estimation. SIAM, Philadelphia
- Tatman C (2001) Production rate selection for steeply dipping tabular deposits. *Min Eng* 53:34–36
- Tong A (1996) Stochastic parameter estimation, reliability analysis, and experimental design in ground water modeling, PhD dissertation. University of California, Los Angeles
- Vallee M (2000) Mineral resources + engineering, economic and legal feasibility = ore reserve. *Bull Can Soc Min Eng* 93:53–61
- Vargaz-Guzman J, Dimitrakopoulos R (2002) Conditional simulation of random fields by successive residuals. *Math Geol* 34:597–611
- Vargaz-Guzman J, Yeh TC (1999) Sequential kriging and cokriging: two powerful geostatistical approaches. *Stoch Environ Res Risk Assess* 13:416–435
- Wilson J, Kitanidis P, Dettinger M (1978) State and parameter estimation in groundwater models. In: Applications of Kalman filter to hydrology, hydraulics and water resources, University of Pittsburgh
- Yates S, Warrick A (1987) Estimating soil water content using co-kriging. *Soil Sci Soc Am J* 51:23–30
- Yeh TC, Zhang J (1996) A geostatistical inverse method for variable saturated flow in the vadose zone. *Water Resour Res* 32:2757–2766
- Zhou H, Gomez-Hernandez J, Hendricks Franssen HJ, Li L (2011) An approach to handling non-gaussianity of parameters and state variables in ensemble kalman filtering. *Adv Water Resour* 34:844–864
- Zimmer B (2012) Continuous mining equipment vs. complex geology challenges in mine planning. In: 11th Inter Symposium of Continuous surface mining, Miskolc, The university of Miskolc