

Overview of WRF Data Assimilation

Tom Auligné

National Center for Atmospheric Research, Boulder, CO USA

WRFDA Tutorial - August 3-5, 2010



Acknowledgments and References

- WRF Tutorial Lectures (H. Huang and D. Barker)
- Data Assimilation concepts and methods
(ECMWF Training Course, F. Bouttier and P. Courtier)
- Data Assimilation Research Testbed (DART) Tutorial
(J. Anderson et al.)
- Analysis methods for numerical weather prediction
(A.C. Lorenc, 1986, *Quart. J. R. Meteorol. Soc.*)
- Atmospheric Data Analysis
(R. Daley, 1991, *Cambridge University Press*, 457 pp.)
- Atmospheric Modeling, Data Assimilation and Predictability
(E. Kalnay, 2003, *Cambridge University Press*, 341 pp.)

Table of contents

- 1 Introduction
- 2 Simple Scalar Example
 - Kalman Filter equations
- 3 Modern Implementations
 - Sequential Algorithms
 - Ensemble Kalman Filter
 - 3D Variational (3DVar)
 - Smoothers
 - 4D Variational (4DVar)
- 4 WRFDA Overview

Motivation

- *A sufficiently accurate knowledge of the state of the atmosphere at the initial time.*

(Today's weather)

- *A sufficiently accurate knowledge of the laws according to which one state of the atmosphere develops from another.*

(Tomorrow's weather)



Vilhelm Bjerknes (1904)
(Peter Lynch)

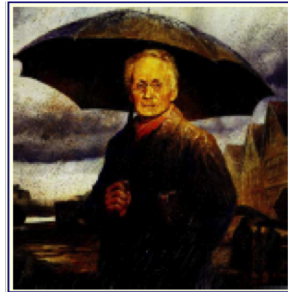
Motivation

- *A sufficiently accurate knowledge of the state of the atmosphere at the initial time.*

(Today's weather)

- *A sufficiently accurate knowledge of the laws according to which one state of the atmosphere develops from another.*

(Tomorrow's weather)



Vilhelm Bjerknes (1904)
(Peter Lynch)

Motivation

- Initial conditions for Numerical Weather Prediction (NWP)
- Calibration and validation
- Observing system design, monitoring and assessment
- Reanalysis
- Better understanding (Model errors, Data errors, Physical process interactions, *etc*)

From Empirical to Statistical methods

- Successive Correction Method (SCM, *Cressman 1959*)
Each observation within a radius of influence L is given a weight w varying with the distance r to the model grid point:
$$w(r) = \frac{L^2 - r^2}{L^2 + r^2} (r \leq L)$$
- Nudging
- Physical Initialization (PI), Latent Heat Nudging (LHN)

However...

- Relaxation functions are somewhat arbitrary
- **Good** forecast can be replaced by **bad** observations
- Noisy observations can create unphysical analysis

So...

Modern DA techniques are usually statistical

From Empirical to Statistical methods

- Successive Correction Method (SCM, *Cressman 1959*)
Each observation within a radius of influence L is given a weight w varying with the distance r to the model grid point:
$$w(r) = \frac{L^2 - r^2}{L^2 + r^2} (r \leq L)$$
- Nudging
- Physical Initialization (PI), Latent Heat Nudging (LHN)

However...

- Relaxation functions are somewhat arbitrary
- **Good** forecast can be replaced by **bad** observations
- Noisy observations can create unphysical analysis

So...

Modern DA techniques are usually statistical

What is the temperature in this room?

Notations

- x_t : "True" state
- x_o : Observation
- x_b : Background information
- $d = x_o - x_b$: Innovation or *Departure*

Hypotheses

- Observation and Background errors are uncorrelated, unbiased, normally distributed, with variance R and B *resp.*
- Analysis x_a is "optimal" in RMSE sense
- Linear Analysis: $x_a = \alpha x_o + \beta x_b = x_b + \alpha(x_o - x_b)$

What is the temperature in this room?

Notations

- x_t : "True" state
- x_o : Observation
- x_b : Background information
- $d = x_o - x_b$: Innovation or *Departure*

Hypotheses

- Observation and Background errors are uncorrelated, unbiased, normally distributed, with variance R and B *resp.*
- Analysis x_a is "optimal" in RMSE sense
- Linear Analysis: $x_a = \alpha x_o + \beta x_b = x_b + \alpha(x_o - x_b)$

Best Linear Unbiased Estimate

The analysis value is $x_a = x_b + \alpha(x_o - x_b)$ and its error variance:

$$A = \overline{(x_a - x_t)(x_a - x_t)} = (1 - \alpha)^2 B + \alpha^2 R$$

$$\frac{\partial A}{\partial \alpha} = 2\alpha(B + R) - 2B$$

$$\frac{\partial A}{\partial \alpha} = 0 \Rightarrow \alpha = \frac{B}{B + R}$$

Best Linear Unbiased Estimate

The analysis value is $x_a = x_b + \alpha(x_o - x_b)$ and its error variance:

$$A = \overline{(x_a - x_t)(x_a - x_t)} = (1 - \alpha)^2 B + \alpha^2 R$$

$$\frac{\partial A}{\partial \alpha} = 2\alpha(B + R) - 2B$$

$$\frac{\partial A}{\partial \alpha} = 0 \Rightarrow \alpha = \frac{B}{B + R}$$

Best Linear Unbiased Estimate (BLUE)

$x_a = x_b + K(x_o - x_b)$ with the definition of the *Kalman Gain*:

$$K = B(B + R)^{-1}$$

and the analysis error variance: $A^{-1} = B^{-1} + R^{-1}$

Statistically, the analysis is better than the observation ($A < R$)
and the background ($A < B$)

Best Linear Unbiased Estimate

The analysis value is $x_a = x_b + \alpha(x_o - x_b)$ and its error variance:

$$A = \overline{(x_a - x_t)(x_a - x_t)} = (1 - \alpha)^2 B + \alpha^2 R$$

$$\frac{\partial A}{\partial \alpha} = 2\alpha(B + R) - 2B$$

$$\frac{\partial A}{\partial \alpha} = 0 \Rightarrow \alpha = \frac{B}{B + R}$$

Best Linear Unbiased Estimate (BLUE)

$x_a = x_b + K(x_o - x_b)$ with the definition of the *Kalman Gain*:

$$K = B(B + R)^{-1}$$

and the analysis error variance: $A^{-1} = B^{-1} + R^{-1}$

Statistically, the analysis is better than the observation ($A < R$)
and the background ($A < B$)

Variational Cost Function

This solution is equivalent to minimizing the cost function:

$$J(x) = \frac{1}{2}(x - x_b)^T B^{-1}(x - x_b) + \frac{1}{2}(x - x_o)^T R^{-1}(x - x_o) = \mathbf{J}_b + \mathbf{J}_o$$

Proof:

$$\nabla J = B^{-1}(x - x_b) + R^{-1}(x - x_o) = 0$$

$$\begin{aligned} \Rightarrow x_a &= x_b + \frac{B}{B + R}(x_o - x_b) \\ &= x_b + K(x_o - x_b) \end{aligned}$$

Variational Cost Function

This solution is equivalent to minimizing the cost function:

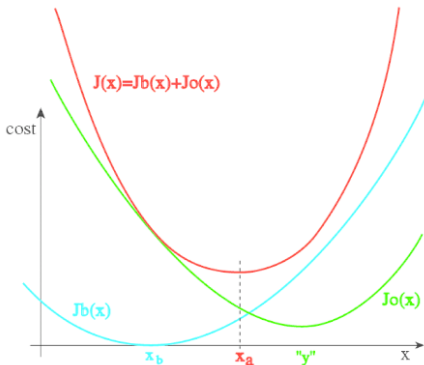
$$J(x) = \frac{1}{2}(x - x_b)^T B^{-1}(x - x_b) + \frac{1}{2}(x - x_o)^T R^{-1}(x - x_o) = \mathbf{J}_b + \mathbf{J}_o$$

Proof:

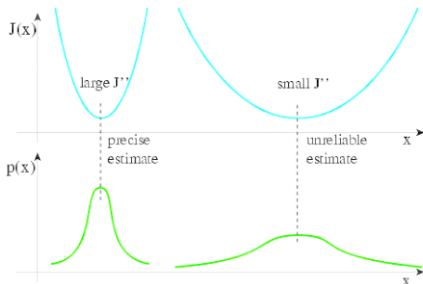
$$\nabla J = B^{-1}(x - x_b) + R^{-1}(x - x_o) = 0$$

$$\begin{aligned}\Rightarrow x_a &= x_b + \frac{B}{B + R}(x_o - x_b) \\ &= x_b + K(x_o - x_b)\end{aligned}$$

Analysis Accuracy



from Bouttier and Courtier 1999



Quality of the Analysis

The precision is defined by the convexity or **Hessian** $A = J''^{-1}$

Conditional Probabilities

According to Bayes Theorem, the joint pdf of x and x_o is:

$$P(x \wedge x_o) = P(x|x_o)P(x_o) = P(x_o|x)P(x)$$

Since $P(x_o) = 1$, $P(x|x_o) = P(x_o|x)P(x)$

We assumed the background and observation errors were Gaussian:

$$P(x) = \lambda_b e^{[\frac{1}{2B}(x_b-x)^2]} \quad \text{and} \quad P(x_o|x) = \lambda_o e^{[\frac{1}{2R}(x_o-x)^2]}$$

$$\Rightarrow P(x|x_o) = \lambda_a e^{[\frac{1}{2R}(x_o-x)^2 + \frac{1}{2B}(x_b-x)^2]} = \lambda_a e^{-J(x)}$$

Maximum Likelihood

The minimum of the cost function J is also the estimator of x_t with the maximum likelihood

Conditional Probabilities

According to Bayes Theorem, the joint pdf of x and x_o is:

$$P(x \wedge x_o) = P(x|x_o)P(x_o) = P(x_o|x)P(x)$$

Since $P(x_o) = 1$, $P(x|x_o) = P(x_o|x)P(x)$

We assumed the background and observation errors were Gaussian:

$$P(x) = \lambda_b e^{[\frac{1}{2B}(x_b-x)^2]} \quad \text{and} \quad P(x_o|x) = \lambda_o e^{[\frac{1}{2R}(x_o-x)^2]}$$

$$\Rightarrow P(x|x_o) = \lambda_a e^{[\frac{1}{2R}(x_o-x)^2 + \frac{1}{2B}(x_b-x)^2]} = \lambda_a e^{-J(x)}$$

Maximum Likelihood

The minimum of the cost function J is also the estimator of x_t with the maximum likelihood

Partial Conclusions

Under the aforementioned hypotheses, the BLUE:

- can be determined analytically through the Kalman gain K
- is also the minimum of a cost function $J = J_b + J_o$
- is optimal for minimum variance **and** maximum likelihood

Sequential Data Assimilation

Forecast model $M_{i \rightarrow i+1} = M$ from step i to $i + 1$

$$x_{i+1}^t = M(x_i^t) + q_i$$

where q_i is the model error. As q_i is unknown and x_i^a is the best estimate of x_i^t , usually: $x_{i+1}^f = M(x_i^a)$

Forecast error

$$x_{i+1}^f - x_{i+1}^t = M(x_i^a) - M(x_i^t) - q_i \approx \mathbf{M}_i(x_i^a - x_i^t) - q_i$$

$\mathbf{M}$ is called the **Tangent-Linear** code of the non-linear model M

Sequential Data Assimilation

Forecast model $M_{i \rightarrow i+1} = M$ from step i to $i + 1$

$$x_{i+1}^t = M(x_i^t) + q_i$$

where q_i is the model error. As q_i is unknown and x_i^a is the best estimate of x_i^t , usually: $x_{i+1}^f = M(x_i^a)$

Forecast error

$$x_{i+1}^f - x_{i+1}^t = M(x_i^a) - M(x_i^t) - q_i \approx \mathbf{M}_i(x_i^a - x_i^t) - q_i$$

$\mathbf{M}$ is called the **Tangent-Linear** code of the non-linear model M

Forecast error covariance matrix

$$P_{i+1}^f \approx \mathbf{M}_i \overline{(x_i^a - x_i^t)(x_i^a - x_i^t)^T} \mathbf{M}_i + \overline{q_i q_i^T} = \mathbf{M}_i P_i^a \mathbf{M}_i^T + Q_i$$

Sequential Data Assimilation

Forecast model $M_{i \rightarrow i+1} = M$ from step i to $i + 1$

$$x_{i+1}^t = M(x_i^t) + q_i$$

where q_i is the model error. As q_i is unknown and x_i^a is the best estimate of x_i^t , usually: $x_{i+1}^f = M(x_i^a)$

Forecast error

$$x_{i+1}^f - x_{i+1}^t = M(x_i^a) - M(x_i^t) - q_i \approx \mathbf{M}_i(x_i^a - x_i^t) - q_i$$

$\mathbf{M}$ is called the **Tangent-Linear** code of the non-linear model M

Forecast error covariance matrix

$$P_{i+1}^f \approx \mathbf{M}_i \overline{(x_i^a - x_i^t)(x_i^a - x_i^t)^T} \mathbf{M}_i + \overline{q_i q_i^T} = \mathbf{M}_i P_i^a \mathbf{M}_i^T + Q_i$$

Sequential Data Assimilation

We can use the forecast as background for the **BLUE** calculation

$$\begin{aligned} K_i &= P_i^f (P_i^f + R)^{-1} \\ x_i^a &= x_i^f + K(x_i^o - x_i^f) \\ (P_i^a)^{-1} &= (P_i^f)^{-1} + R^{-1} \Rightarrow P_i^a = (I - K_i)P_i^f \end{aligned}$$

Finally, we can distinguish the model space x from the observation space y and introduce an Observation Operator $H : x \mapsto y$, which is linearized: $H(x_i^a) - H(x_i^f) \approx \mathbf{H}(x_i^a - x_i^f)$

$$\begin{aligned} K_i &= P_i^f \mathbf{H}_i^T (\mathbf{H}_i P_i^f \mathbf{H}_i^T + R)^{-1} \\ x_i^a &= x_i^f + K(y_i^o - x_i^f) \\ P_i^a &= (I - K_i \mathbf{H}_i) P_i^f \end{aligned}$$

Sequential Data Assimilation

We can use the forecast as background for the **BLUE** calculation

$$\begin{aligned} K_i &= P_i^f (P_i^f + R)^{-1} \\ x_i^a &= x_i^f + K(x_i^o - x_i^f) \\ (P_i^a)^{-1} &= (P_i^f)^{-1} + R^{-1} \Rightarrow P_i^a = (I - K_i)P_i^f \end{aligned}$$

Finally, we can distinguish the model space x from the observation space y and introduce an Observation Operator $H : x \mapsto y$, which is linearized: $H(x_i^a) - H(x_i^f) \approx \mathbf{H}(x_i^a - x_i^f)$

$$\begin{aligned} K_i &= P_i^f \mathbf{H}_i^T (\mathbf{H}_i P_i^f \mathbf{H}_i^T + R)^{-1} \\ x_i^a &= x_i^f + K(y_i^o - x_i^f) \\ P_i^a &= (I - K_i \mathbf{H}_i) P_i^f \end{aligned}$$

The Extended Kalman Filter Algorithm

Analysis step i :

$$K_i = P_i^f \mathbf{H}_i^T [\mathbf{H}_i P_i^f \mathbf{H}_i^T + R]^{-1} \quad (1)$$

$$\mathbf{x}_i^a = \mathbf{x}_i^f + K_i [y^o - H \mathbf{x}_i^f] \quad (2)$$

$$P_i^a = [I - K_i \mathbf{H}_i] P_i^f \quad (3)$$

Forecast step from i to $i + 1$:

$$\mathbf{x}_{i+1}^f = M(\mathbf{x}_i^a) \quad (4)$$

$$P_{i+1}^f = \mathbf{M}_i P_i^a \mathbf{M}_i^T + Q_i \quad (5)$$

Hypotheses

- Gaussian distributions of errors
- **M**: Linearization around non-linear Model M
- **H**: Linearization around non-linear Observation Operator H

The Extended Kalman Filter Algorithm

Analysis step i :

$$K_i = P_i^f \mathbf{H}_i^T [\mathbf{H}_i P_i^f \mathbf{H}_i^T + R]^{-1} \quad (1)$$

$$\mathbf{x}_i^a = \mathbf{x}_i^f + K_i [y^o - H \mathbf{x}_i^f] \quad (2)$$

$$P_i^a = [I - K_i \mathbf{H}_i] P_i^f \quad (3)$$

Forecast step from i to $i + 1$:

$$\mathbf{x}_{i+1}^f = M(\mathbf{x}_i^a) \quad (4)$$

$$P_{i+1}^f = \mathbf{M}_i P_i^a \mathbf{M}_i^T + Q_i \quad (5)$$

Hypotheses

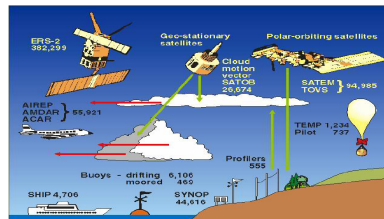
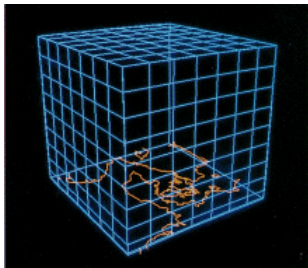
- Gaussian distributions of errors
- **M**: Linearization around non-linear Model M
- **H**: Linearization around non-linear Observation Operator H

From scalar to vector: dimensions

$$x \rightarrow \mathbf{x}$$

Number of grid points $\approx 10^7$

Dimension of P^f , $P^a \approx 10^7 \times 10^7$



$$y^o \rightarrow \mathbf{y}^o$$

Number of observations $\approx 10^6$

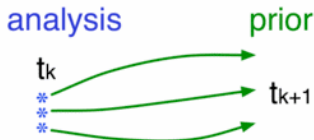
Dimension of $R \approx 10^6 \times 10^6$

Ensemble Kalman Filter (EnKF)

Hypotheses

- Monte Carlo approximation to pdfs
- Gaussian distributions used for computing update
- Localization in space: for each model grid point, only a few observations are used to compute the analysis increment.

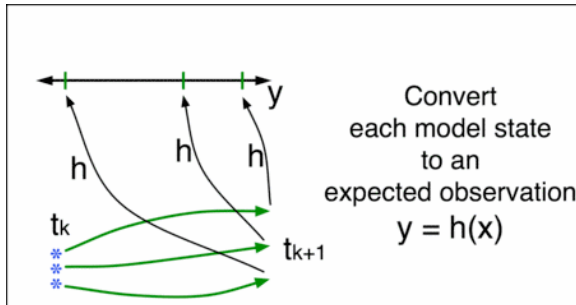
3 ensemble members advancing in time



Ensemble Kalman Filter (EnKF)

Hypotheses

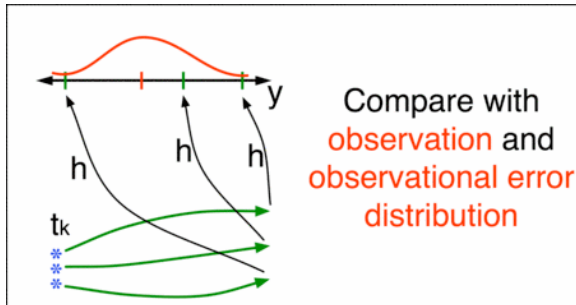
- Monte Carlo approximation to pdfs
- Gaussian distributions used for computing update
- Localization in space: for each model grid point, only a few observations are used to compute the analysis increment.



Ensemble Kalman Filter (EnKF)

Hypotheses

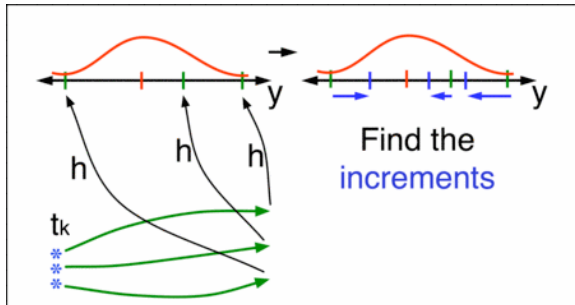
- Monte Carlo approximation to pdfs
- Gaussian distributions used for computing update
- Localization in space: for each model grid point, only a few observations are used to compute the analysis increment.



Ensemble Kalman Filter (EnKF)

Hypotheses

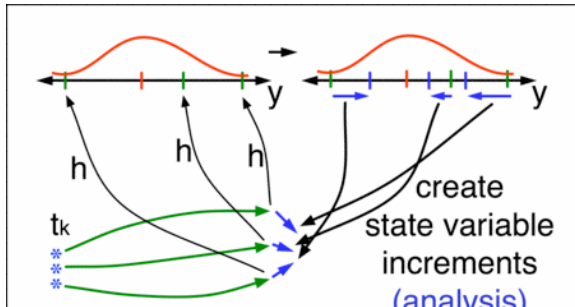
- Monte Carlo approximation to pdfs
- Gaussian distributions used for computing update
- Localization in space: for each model grid point, only a few observations are used to compute the analysis increment.



Ensemble Kalman Filter (EnKF)

Hypotheses

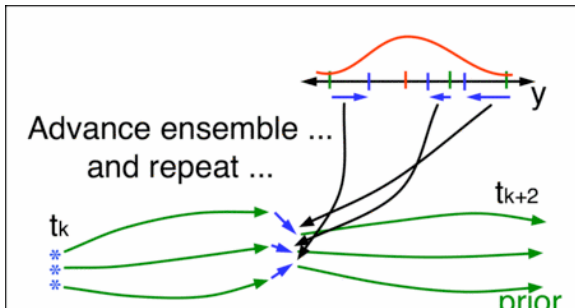
- Monte Carlo approximation to pdfs
- Gaussian distributions used for computing update
- Localization in space: for each model grid point, only a few observations are used to compute the analysis increment.



Ensemble Kalman Filter (EnKF)

Hypotheses

- Monte Carlo approximation to pdfs
- Gaussian distributions used for computing update
- Localization in space: for each model grid point, only a few observations are used to compute the analysis increment.



Ensemble Kalman Filter (EnKF)

Hypotheses

- Monte Carlo approximation to pdfs
- Gaussian distributions used for computing update
- Localization in space: for each model grid point, only a few observations are used to compute the analysis increment.

Advantages

- Easy to implement and provides estimate of Analysis Accuracy
- H and M need not be linearized

Drawbacks

Localization avoids degeneracy from under-sampling and reduces spurious noise, but it affects model internal balance

3D Variational Data Assimilation (3DVar)

Hypotheses

Avoid calculating K by solving the equivalent minimization problem defined by the cost function:

$$J(x) = \frac{1}{2}(x - x_b)^T B^{-1}(x - x_b) + \frac{1}{2}(y^o - H(x))^T R^{-1}(y^o - H(x))$$

$$\nabla J(x) = B^{-1}(x - x_b) - \mathbf{H}^T R^{-1}[y - H(x)]$$

$\mathbf{H}^T$ is called the **Adjoint** of the linearized observation operator

3D Variational Data Assimilation (3DVar)

Hypotheses

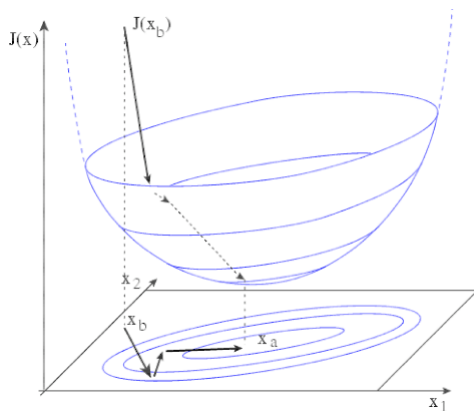
Avoid calculating K by solving the equivalent minimization problem defined by the cost function:

$$J(x) = \frac{1}{2}(x - x_b)^T B^{-1}(x - x_b) + \frac{1}{2}(y^o - H(x))^T R^{-1}(y^o - H(x))$$

$$\nabla J(x) = B^{-1}(x - x_b) - \mathbf{H}^T R^{-1}[y - H(x)]$$

$\mathbf{H}^T$ is called the **Adjoint** of the linearized observation operator

3D Variational Data Assimilation (3DVar)



from Bouttier and Courtier 1999

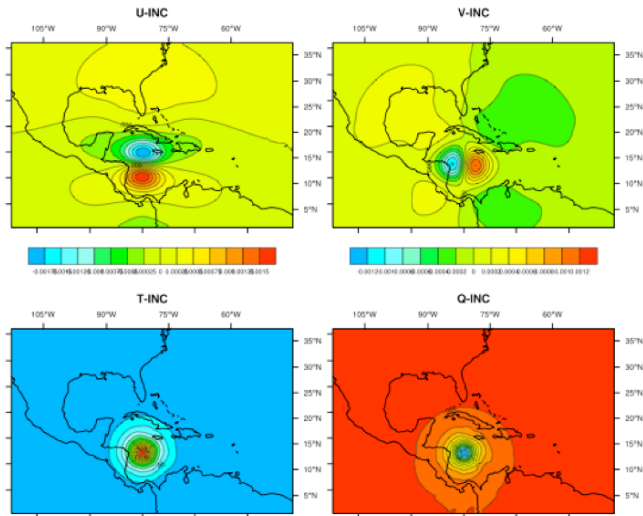
Minimization Algorithm

- Iterative minimizer
→ several simulations
- Steepest Descent, Quasi-Newton, Conjugate Gradient, *etc*

Preconditioning

- Improve Condition Nb
- Faster convergence

Single Observation Experiment



3D Variational Data Assimilation (3DVar)

Hypotheses

- Avoid calculating K by solving the equivalent minimization problem defined by the cost function:

$$J(x) = \frac{1}{2}(x - x_b)^T B^{-1}(x - x_b) + \frac{1}{2}(y^o - H(x))^T R^{-1}(y^o - H(x))$$

Advantages

- Easy to use with complex observation operators
- Can add external weak or *penalty* constraints J_c

Drawbacks

- Sub-optimal for strongly non-linear observation operators
- All observations are assumed to be instantaneous

3D Variational Data Assimilation (3DVar)

Hypotheses

- Avoid calculating K by solving the equivalent minimization problem defined by the cost function:

$$J(x) = \frac{1}{2}(x - x_b)^T B^{-1}(x - x_b) + \frac{1}{2}(y^o - H(x))^T R^{-1}(y^o - H(x))$$

Advantages

- Easy to use with complex observation operators
- Can add external weak or *penalty* constraints J_c

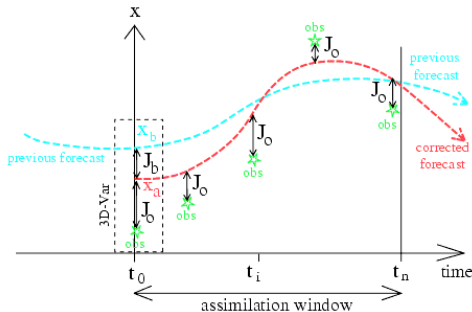
Drawbacks

- Sub-optimal for strongly non-linear observation operators
- All observations are assumed to be instantaneous

4D Variational Data Assimilation (4DVar)

Hypotheses

- Generalization of 3DVar for observations distributed in time
- Analysis variable x defined at the **beginning** of time window
- Find model trajectory minimizing the distance to observations



4D Variational Data Assimilation (4DVar)

Hypotheses

- Generalization of 3DVar for observations distributed in time
- Analysis variable x defined at the **beginning** of time window
- Find model trajectory minimizing the distance to observations

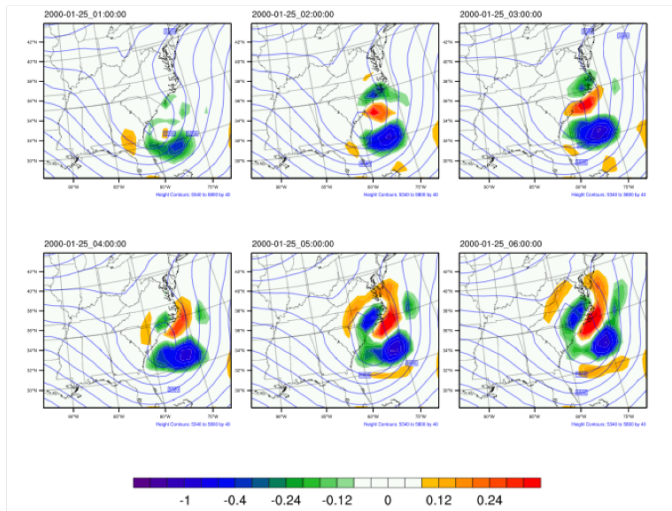
The Cost Function becomes:

$$J(x) = \frac{1}{2}(x - x_b)^T B^{-1}(x - x_b) + \frac{1}{2}(y^o - HM(x))^T R^{-1}(y^o - HM(x))$$

$$\nabla J(x) = B^{-1}(x - x_b) - \mathbf{M}^T \mathbf{H}^T R^{-1}[y - HM(x)]$$

$\mathbf{M}^T$ is called the **Adjoint** of the linearized forecast model

4D Variational Data Assimilation (4DVar)



4D Variational Data Assimilation (4DVar)

Hypotheses

- Generalization of 3DVar for observations distributed in time
- Analysis variable x defined at the **beginning** of time window
- Find model trajectory minimizing the distance to observations

Advantages

Model internal balance is more prone to be respected

Drawbacks

- The development and maintenance of the Adjoint model $\mathbf{M}^T$ can be cumbersome
- Limitation of the "perfect model" assumption

4D Variational Data Assimilation (4DVar)

Hypotheses

- Generalization of 3DVar for observations distributed in time
- Analysis variable x defined at the **beginning** of time window
- Find model trajectory minimizing the distance to observations

Advantages

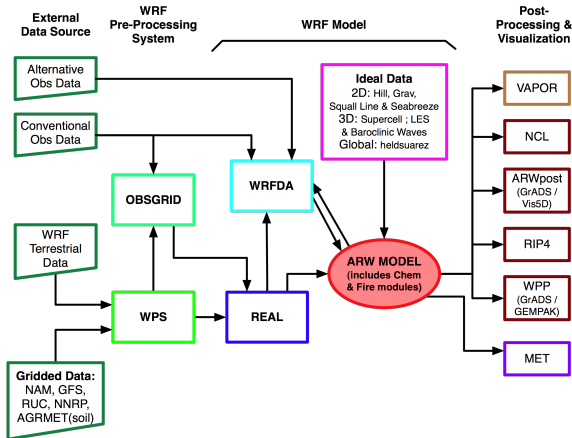
Model internal balance is more prone to be respected

Drawbacks

- The development and maintenance of the Adjoint model $\mathbf{M}^T$ can be cumbersome
- Limitation of the "perfect model" assumption

WRF Data Assimilation (WRFDA)

WRF Modeling System Flow Chart



WRF Data Assimilation (WRFDA)

Community WRF DA System

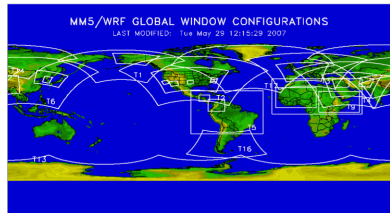
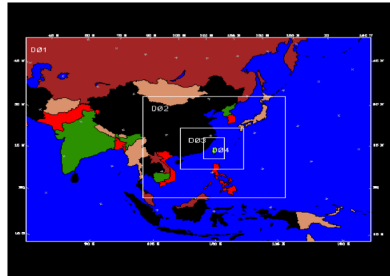
- Regional/Global
- Research/Operations
- Deterministic/Probabilistic

Algorithms

- 3DVar, 4DVar (Regional)
- Ensemble (ETKF/EnKF)
- Hybrid Var/Ens

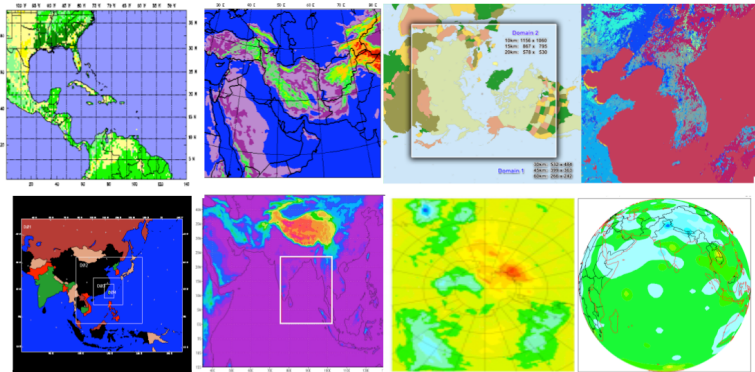
Model: WRF

ARW, NMM



WRFDA Program

- NCAR Staff: 20FTE, 10 projects
- Ext. collaborators (AFWA, KMA, CWB, BMB): 10 FTE
- Community Users: 40



WRFDA Observations

Conventional

- Surface (SYNOP, METAR, SHIP, BUOY)
- Upper Air (TEMP, PIBAL, AIREP, ACARS, TAMDAR)

Bogus

- Tropical Cyclone Bogus
- Global Bogus

WRFDA Observations

Remotely Sensed Retrievals

- Atmospheric Motion Vectors (from GEOs and Polar)
- SATEM Thickness
- Ground-based GPS TPW/Zenith Total Delay
- SSM/I oceanic surface wind speed and TPW
- Scatterometer oceanic surface winds
- Wind Profiler
- Radar Radial Velocities and Reflectivities
- Satellite Temperature, humidity, thickness profiles
- GPS Refractivity (COSMIC)

WRFDA Observations

Satellite Radiances (RTTOV or CRTM Radiative Transfer)

- HIRS (from NOAA-16, 17, 18 and METOP-2)
- AMSU-A (from NOAA-15, 16, 18, EOS-Aqua and METOP-2)
- AMSU-B (from NOAA-15, 16, 17)
- MHS (from NOAA-18 and METOP-2)
- AIRS (from EOS-Aqua)
- SSMIS (from DMSP-16)

www.mmm.ucar.edu/wrf/users/wrfda

The screenshot shows a web browser window titled "WRFDA Model Users Site". The address bar shows the URL "http://www.mmm.ucar.edu/wrf/users/wrfda/". The browser's "Most Visited" list includes "Gmail - Inbox - hms.xy.huang", "Google Calendar", and "WRFDA Model Users Site". The website has a green header with the text "WRFDA USERS PAGE" and a background image of a weather map. Below the header is a navigation bar with links: "Home", "Analysis System", "User Support", "Download", "Doc / Pub", "Links", and "Users Forum". A search bar is located on the right side of the navigation bar. The main content area is divided into two columns. The left column is green and contains links: "wrf-model.org", "Public Domain Notice", and "Contact WRF Support". The right column is white and contains the following content: a section header "WRF Data Assimilation System Users Page" in red, a welcome message, an "ANNOUNCEMENTS" section with links to "WRF Tutorials", "WRF Version 3.1 Release Information", "WRF Version 3.0.1.1 Release", and "WRF Var Version 3.0.1.1 Release", and a notice about the 9th WRF Users' Workshop.

WRFDA Model Users Site

http://www.mmm.ucar.edu/wrf/users/wrfda/

Most Visited -

Gmail - Inbox - hms.xy.huang Google Calendar WRFDA Model Users Site

WRFDA USERS PAGE

Home Analysis System User Support Download Doc / Pub Links Users Forum

Search

wrf-model.org

Public Domain Notice

Contact WRF Support

WRF Data Assimilation System Users Page

Welcome to the users home page for the Weather Research and Forecasting (WRF) model data assimilation system (WRFDA). The WRFDA system is in the public domain and is freely available for community use. It is designed to be a flexible, state-of-the-art atmospheric data assimilation system that is portable and efficient on available parallel computing platforms. WRFDA is suitable for use in a broad range of applications across scales ranging from kilometers of regional mesoscale to thousands of kilometers of global scales.

The Mesoscale and Microscale Meteorology Division of NCAR is currently maintaining and supporting a subset of the overall WRF code (Version 3) that includes:

ANNOUNCEMENTS

[WRF Tutorials](#) - January 26 - February 5, 2009, Boulder, Colorado.

[WRF Version 3.1 Release Information](#)

[WRF Version 3.0.1.1 Release:](#)
August 22, 2008

[WRF Var Version 3.0.1.1 Release:](#)
August 29, 2008

New 'Known Problems' posts for V3
[WRF](#) (1/6/09) and [WPS](#) (8/4/08)

The 9th WRF Users' Workshop was held June 23 - 27, 2008 in Boulder, Colorado. [Workshop Presentations](#) is now online.

Conclusions

- Observations y^o
- Background x_b
- Observation Operator H
- Innovations $y^o - H(x_b)$

- Observation Error R
- Bkg/Ana Error P^f, P^a
- Tangent-Linear $\mathbf{H}, \mathbf{M}$
- Adjoint $\mathbf{H}^T, \mathbf{M}^T$

(Extended) Kalman Filter (quasi-)linear statistical algorithm

Conclusions

- Observations y^o
- Background x_b
- Observation Operator H
- Innovations $y^o - H(x_b)$

- Observation Error R
- Bkg/Ana Error P^f, P^a
- Tangent-Linear $\mathbf{H}, \mathbf{M}$
- Adjoint $\mathbf{H}^T, \mathbf{M}^T$

(Extended) Kalman Filter (quasi-)linear statistical algorithm

Simplifications for practical implementation

- Ensemble methods: EnKF
- Variational methods: 3DVar, 4DVar

Conclusions

- Observations y^o
- Background x_b
- Observation Operator H
- Innovations $y^o - H(x_b)$

- Observation Error R
- Bkg/Ana Error P^f, P^a
- Tangent-Linear $\mathbf{H}, \mathbf{M}$
- Adjoint $\mathbf{H}^T, \mathbf{M}^T$

(Extended) Kalman Filter (quasi-)linear statistical algorithm

Simplifications for practical implementation

- Ensemble methods: EnKF
- Variational methods: 3DVar, 4DVar

Conclusions

Warning

WRFDA should NOT be used as a *black box*

- Processing of Observations (Quality Control, Bias Correction)
- Modeling of Background and Observation error covariances
- Accounting for Model errors and Non-Linearities

Thank you for your attention...