



Búsqueda de fuentes puntuales con el algoritmo EM en el telescopio de Neutrinos ANTARES



**XXX Reunión Bienal Real Sociedad
Española de Física**

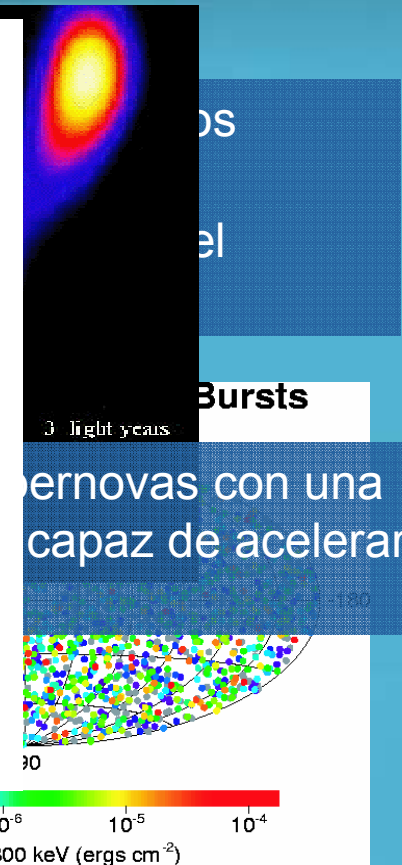
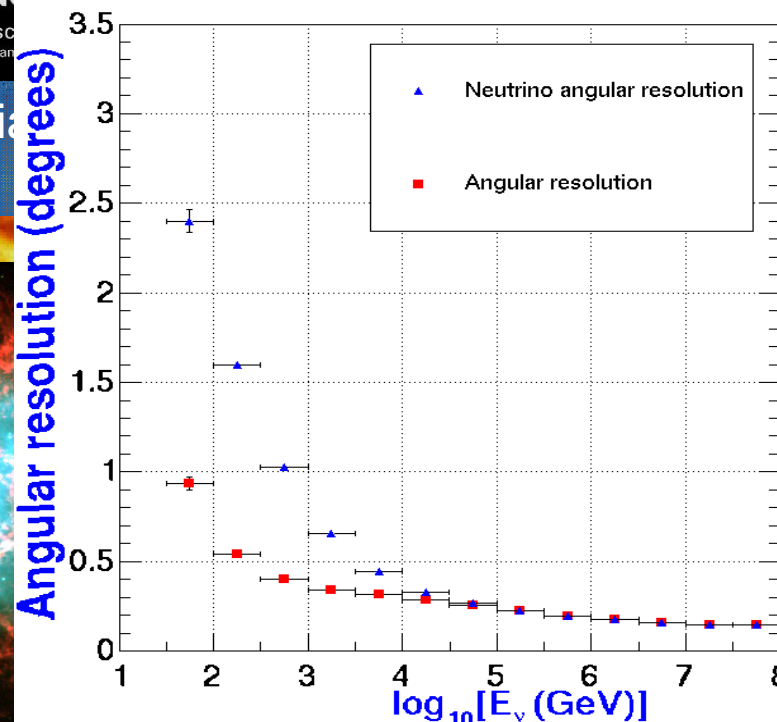
Física de Altas Energías

Juan Antonio Aguilar Sánchez, IFIC

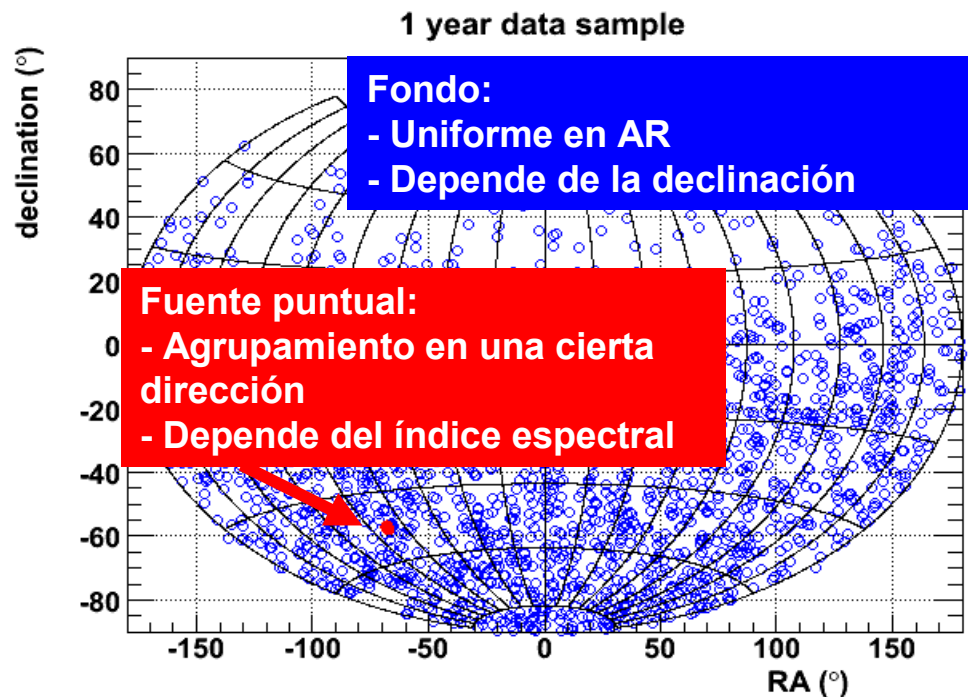
Motivación

- La detección de fuentes emisoras de neutrinos sería un evidencia de los modelos de aceleración hadrónica en los procesos de aceleración de rayos cósmicos (mecanismo Fermi).
- ANTARES** tiene una gran resolución angular por lo que tiene un gran potencial para la búsqueda de fuentes de neutrinos.

Microcuásares: Binarias
radio jets relativistas.



Búsqueda de fuentes puntuales



Algoritmo EM

- Búsqueda de clusters.
- Fuentes con distribuciones Gausianas sobre una distribución de fondo.
- Significancia basada en probabilidad del fondo para producir acumulación de sucesos.

- Neutrinos atmosféricos constituyen el mayor fondo en un telescopio de neutrinos.

~1832 atmospheric ν / year

+ 138 single μ / year

+ 98 multi- μ / year

= ~ 2068 bg. events

- El flujo de neutrinos cósmicos es muy pequeño, pero se concentran en ciertas direcciones.
- Búsqueda de fuentes puntuales: **identificación de agrupamientos de sucesos sobre un fondo de neutrinos atmosféricos.**

Algoritmo EM: Modelo de Mezcla

Construimos la pdf basándonos en modelo de mezclas:

$$p(\mathbf{x}) = \sum_{j=1}^g \pi_j p(\mathbf{x}; \theta_j)$$

- g es el número de componentes en el modelo
- $\pi_j \geq 0$, proporciones de mezcla ($\sum \pi_j = 1$)
- $p(\mathbf{x}; \theta_j)$, $j=1, \dots, g$ funciones densidad de las componentes

Nuestro caso: 1) Pdf = fondo + señal

2) $\mathbf{x} = (\alpha_{RA}, \delta)$

$$\sum_{j=1}^g \pi_j p(\mathbf{x}; \theta_j) = \underbrace{\pi_1 \frac{1}{2\pi} P_7(\delta)}_{\text{Fondo}} + \underbrace{\sum_{j=2}^g \pi_j p_{Gauss}(\alpha_{RA}, \delta | \mu_j, \Sigma_j)}_{\text{Fuentes}}$$

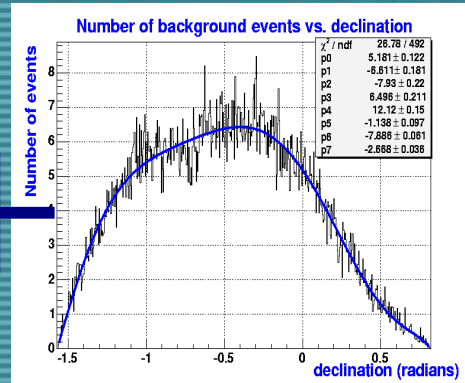
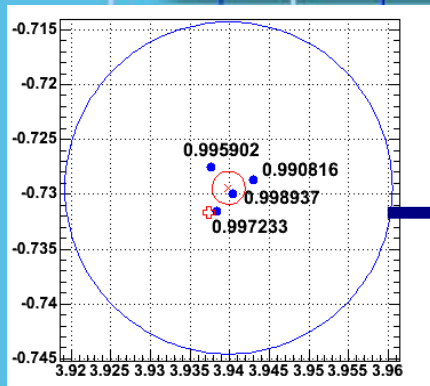
No se usa la energía de los sucesos!

Asumimos que las distribuciones de densidad de las fuentes son Gaussianas:

$$p_{Gauss}(\mathbf{x}; \mu_j, \Sigma_j) = \frac{\exp\left\{-\frac{1}{2}(\mathbf{x} - \mu_j)^T \Sigma_j^{-1}(\mathbf{x} - \mu_j)\right\}}{\sqrt{\det(2\pi \Sigma_j)}}$$

$$\theta_j = (\mu_j, \Sigma_j) \text{ donde } \mu = (\mu_x, \mu_y) \quad \Sigma = (\sigma_x, \sigma_y, \sigma_{xy})$$

Posición de la fuente





Algoritmo EM: Método General

Dado un conjunto de n observaciones la verosimilitud es:

$$L(\Psi) = p(\{\mathbf{x}\}, \Psi) = \prod_{i=1}^n \sum_{j=1}^g \pi_j p(\mathbf{x}_i; \theta_j)$$

- Donde Ψ es un vector que denota el conjunto de parámetros $\{\pi_1, \dots, \pi_g, \theta_1, \dots, \theta_g\}$
- Esta verosimilitud en principio no es maximizable analíticamente!

Conjunto INCOMPLETO

$$\{\mathbf{x}\} = (\alpha_{RA}, \delta)$$

Conjunto COMPLETO

$$\{\mathbf{y}\} \quad y_i = (\mathbf{x}_i, \mathbf{z}_i)$$

$$z_{ik} \begin{cases} 1 & \text{si } \mathbf{x}_i \text{ pertenece al grupo } k \\ 0 & \text{otros} \end{cases}$$

Las nuevas funciones de densidad: $g(\mathbf{y}; \Psi) = g(\mathbf{x}, \mathbf{z}; \Psi) = f(\mathbf{x} | \mathbf{z}, \Psi) p(\mathbf{x}; \theta)$

Verosimilitud del conjunto completo: $L'(\Psi) = g(\{\mathbf{y}\}, \Psi)$

La verosimilitud “incompleta” se obtiene integrando sobre los valores posibles de $\{\mathbf{y}\}$ donde $\{\mathbf{x}\}$ está embebido:

$$L(\Psi) = p(\{\mathbf{x}\}, \Psi) = \int \prod_{i=1}^n g(\mathbf{x}_i, \mathbf{z}; \Psi) d\mathbf{z}$$



E-Step and M-Step

Dos pasos: Expectation-step y Maximization-step:

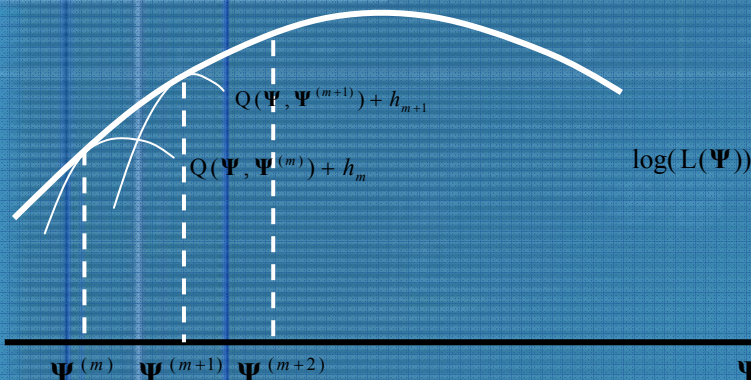
□ E-Step:

- Empezamos con un conjunto inicial de parámetros $\Psi^{(m)}$
- Calculamos el valor esperado de la log-likelihood del conjunto completo, condicionado en los datos observados $\{\mathbf{x}\}$

$$Q(\Psi, \Psi^{(m)}) = E[\log(g(\{\mathbf{y}\}; \Psi)) \mid p(\{\mathbf{x}\}; \Psi^{(m)})]$$

□ M-Step:

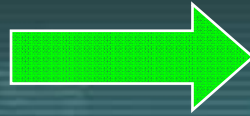
- Encontrar $\Psi = \{\Psi^{(m+1)}\}$ que maximiza $Q(\Psi, \Psi^{(m)})$



Maximizaciones sucesivas de la función $Q(\Psi, \Psi^{(m)})$ llevan a un incremento en la log-likelihood

Selección de modelo (BIC)

1. Conjunto de datos D
2. Varios modelos $M_1, \dots, M_k$



1. M_0 = fondo
2. M_1 = fondo + señal

Necesitamos un parámetro $\lambda(D)$ que nos diga si nuestros datos se ajustan mejor al modelo M_0 o el modelo M_1

Probabilidad de darse M_k dado D

Integración sobre el espacio de parámetros desconocidos θ_k

$$p(D | M_k) = \int p(D | \theta_k, M_k) p(\theta_k | M_k) d\theta_k$$

Bayesian Information Criterion o BIC*:

g.d.l (parámetros a ajustar)

$$2\log[p(D | M_k)] \approx 2\log p(D | \hat{\theta}_k, M_k) - v_k \log(n) = \text{BIC}_k$$

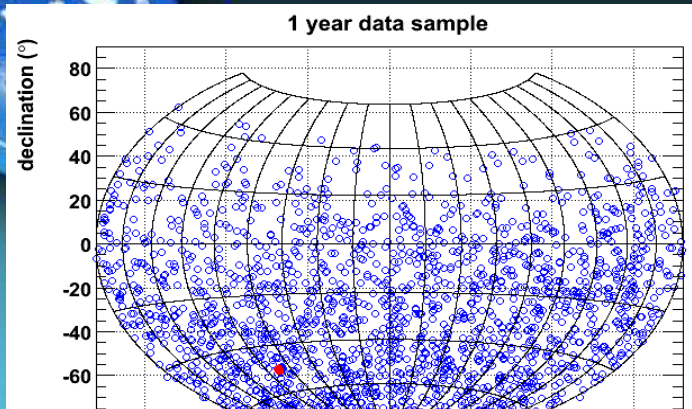
$$\text{BIC}_k = 2 \sum_{i=1}^n \log[p_{\text{Gauss}}(\alpha_{RAi}, \delta_i) + p_{\text{BG}}(\delta_i)] - 2 \sum_{i=1}^n \log[p_{\text{BG}}(\delta_i)] - 6 \log(n)$$

Fondo + señal

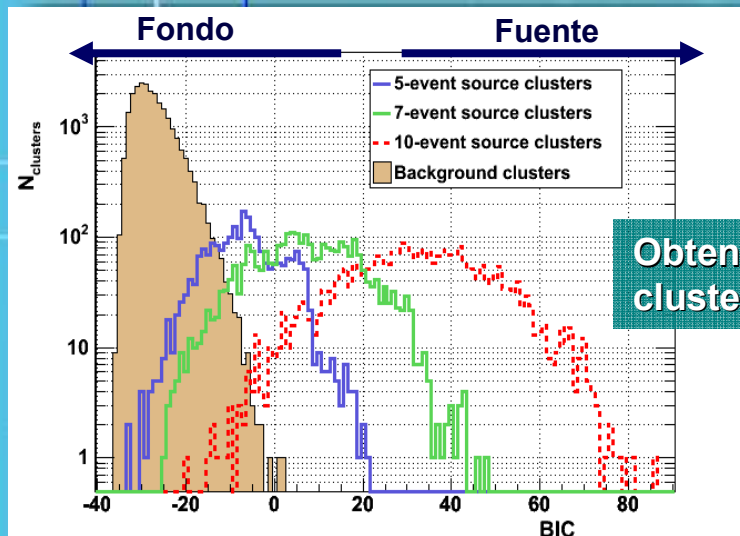
Fondo

6 parámetros

Procedimiento



Realizamos 3000 experimentos cada uno equivalente a 1 año de toma de datos



Obtenemos un valor del BIC por cada cluster candidato de la muestra

Algoritmo sencillo que identifica cada *cluster* de la muestra

Cada cluster se ajusta por el algoritmo EM

Algoritmo EM:

Expectation

Maximization

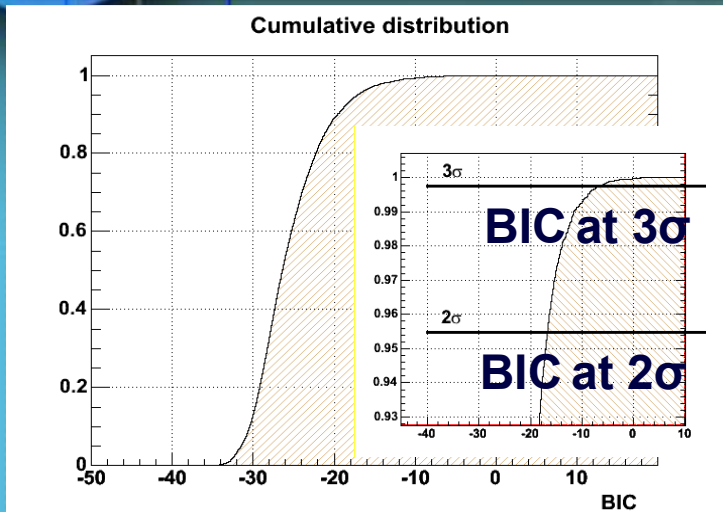
¿Converge?

SI

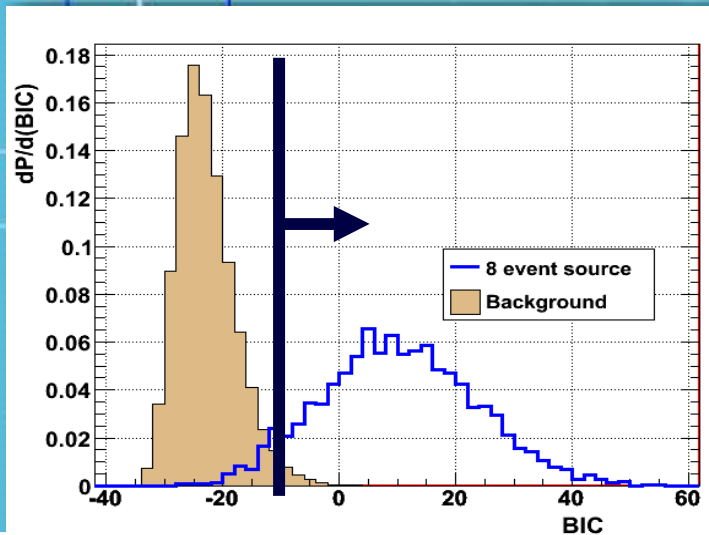
NO



Probabilidad de descubrimiento



- Se calcula la función acumulativa de la distribución de BIC en el caso de sólo fondo
- Valores del BIC para distintos niveles de confianza.



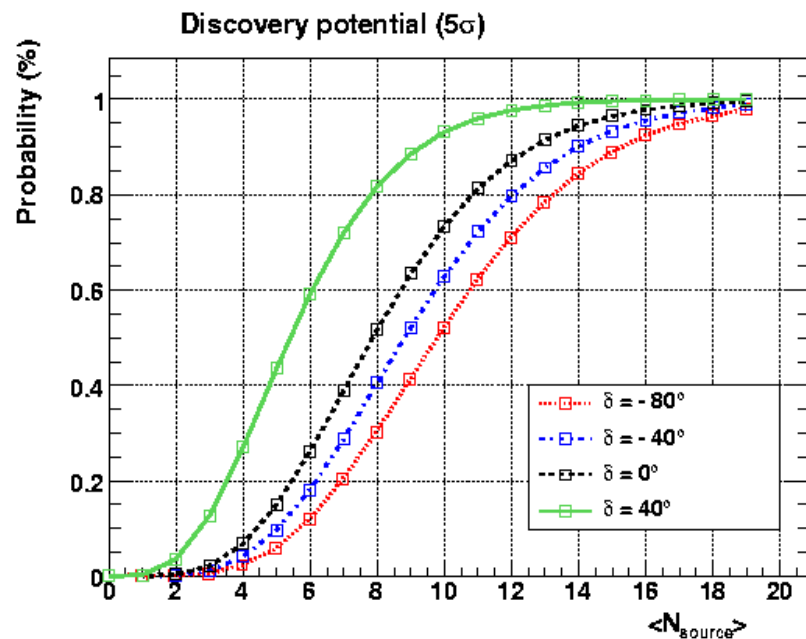
Definición de la probabilidad de descubrimiento es:

$$P(N\sigma) = \frac{\int_{\text{BIC}^{N\sigma}}^{\infty} f_{src}(\text{BIC})}{N_{exp}}$$

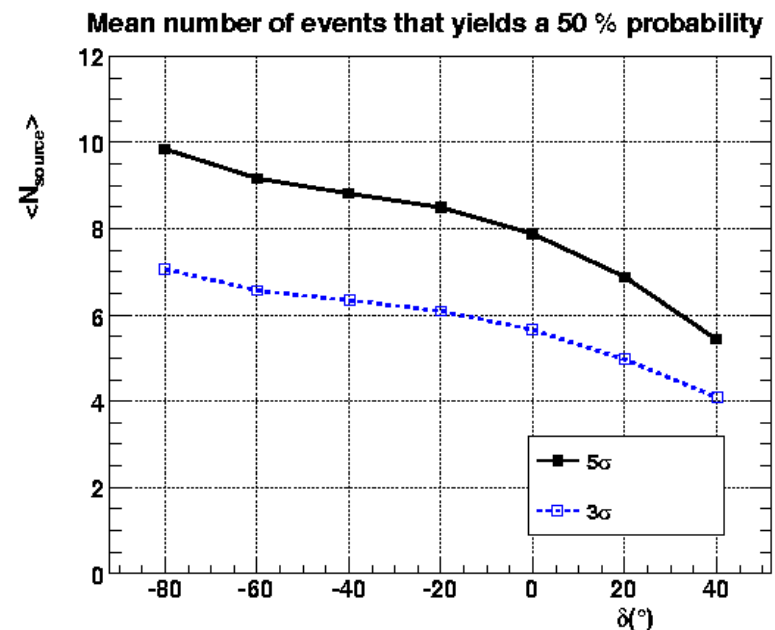
Donde $N\sigma = 2\sigma, 3\sigma, 5\sigma$



Probabilidad de descubrimiento: Resultados



Probabilidad en función del número de sucesos para distintas declinaciones (5σ)

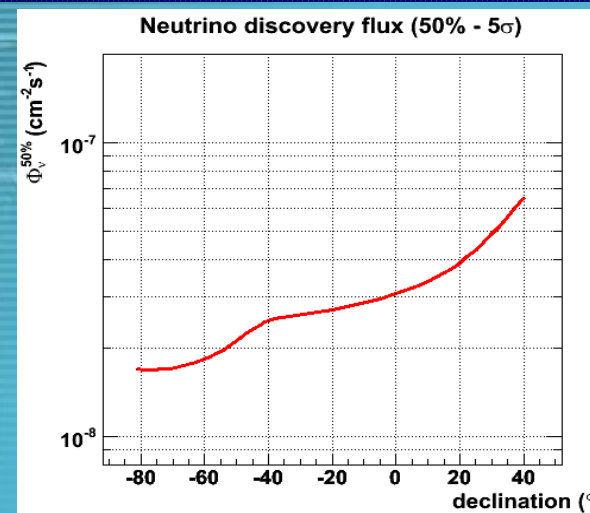
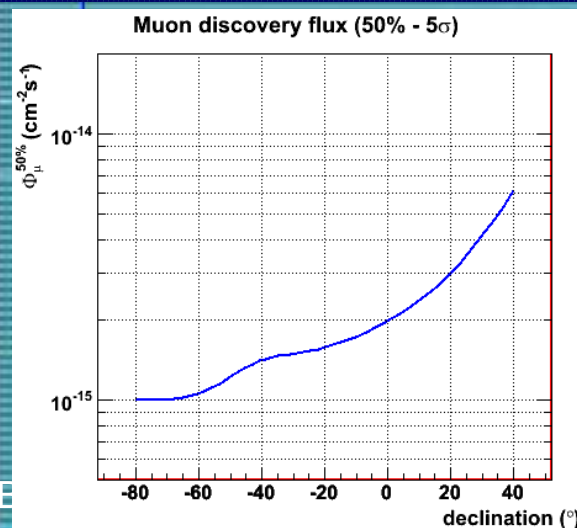
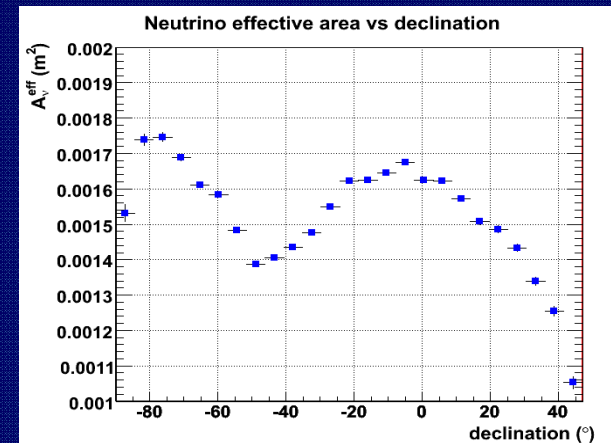
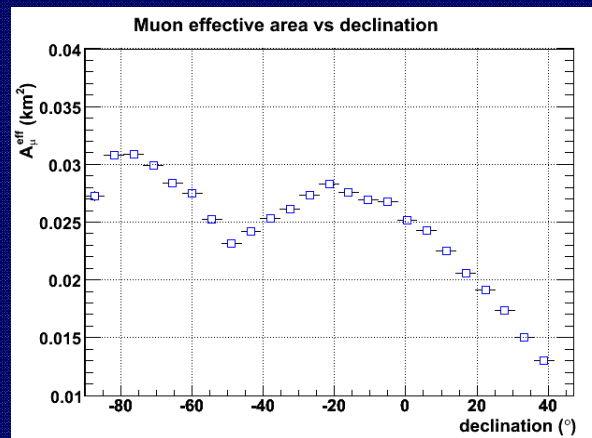
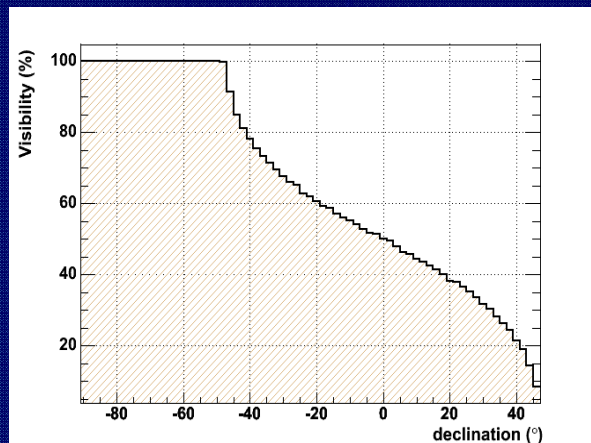


Número de sucesos emitidos por la fuente que dan lugar a una probabilidad del 50% (3σ y 5σ)

Flujo de descubrimiento al 50%

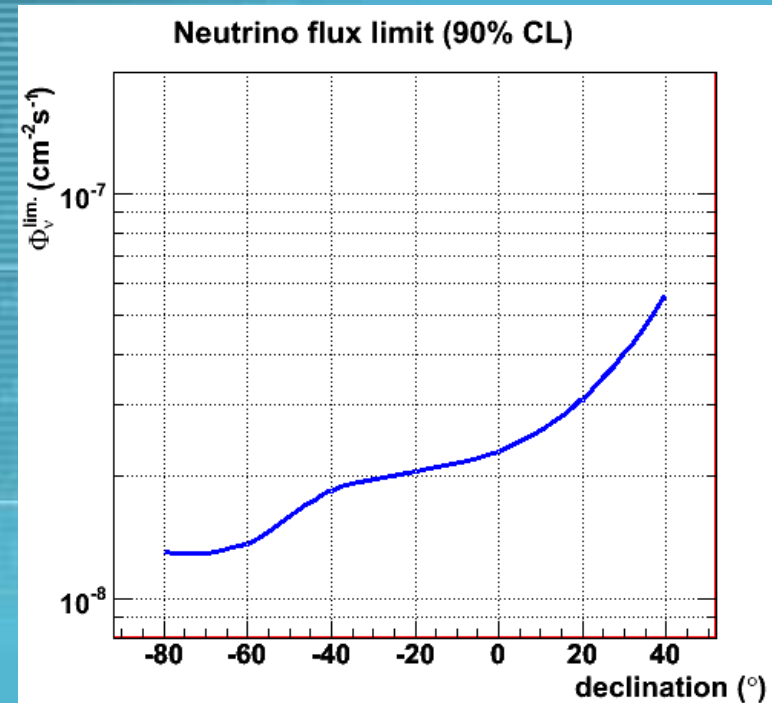
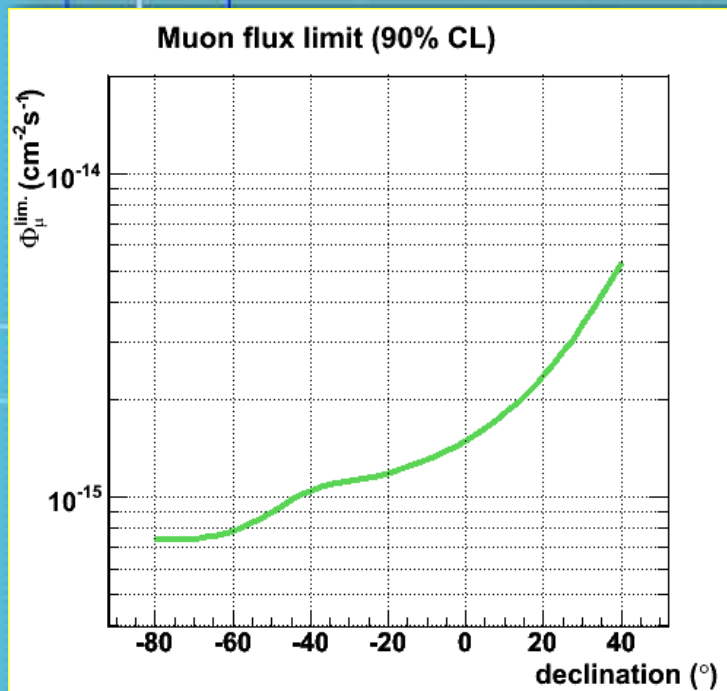
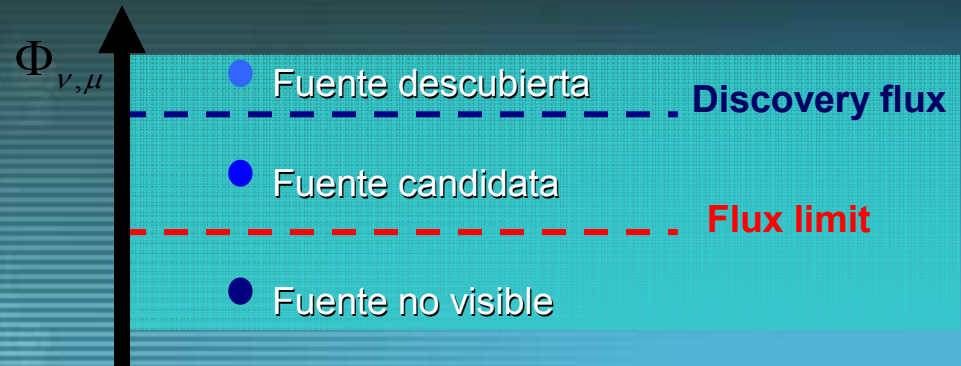
$$\Phi_i(\gamma, \delta) = \frac{N_{events}(\delta)}{A_i^{eff}(\gamma, \delta) V(\delta) T_{live}}$$

- $N_{events}(\delta)$ Número de sucesos al 50%
- $A_i^{eff}(\gamma, \delta)$ Area efectiva ($i = \mu, \nu$) para $\gamma = 2$
- $V(\delta)$ Visibilidad
- T_{live} Tiempo de adquisición del detector



Sensibilidad (flux limit)

Si no hemos visto ninguna fuente, podemos ser capaces de dar un flujo por debajo del cual no podemos decir si existe o no una fuente





Conclusiones

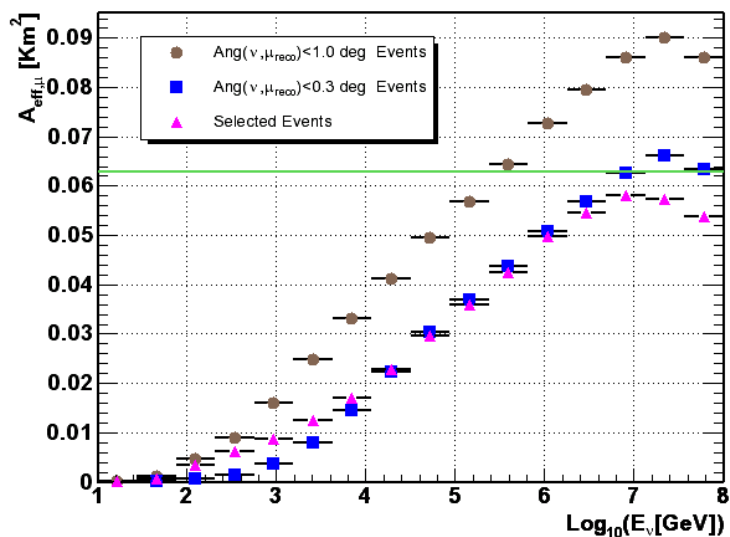
- Se ha estudiado el potencial de **ANTARES** para la búsqueda de fuentes puntuales usando para ello un algoritmo basado en la búsqueda de *clusters*.
- El método basado en el **algoritmo EM** presenta un mayor potencial de descubrimiento que los métodos basado en bins sin la necesidad de usar información estimada del Monte Carlo como otros métodos de búsqueda.
- El **flujo de descubrimiento** al 50% (5σ):

$$\Phi_{\mu}(\delta = 0) \approx 2 \cdot 10^{-15} \text{ cm}^{-2} \text{ s}^{-1} \quad (1 \text{ año})$$

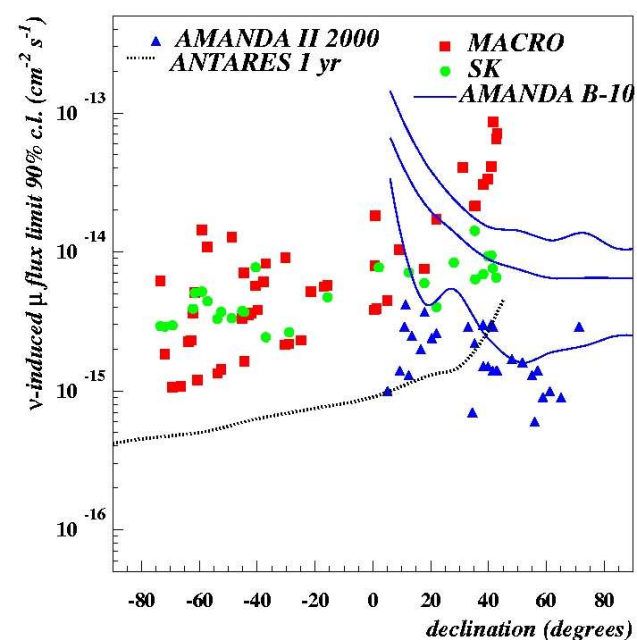
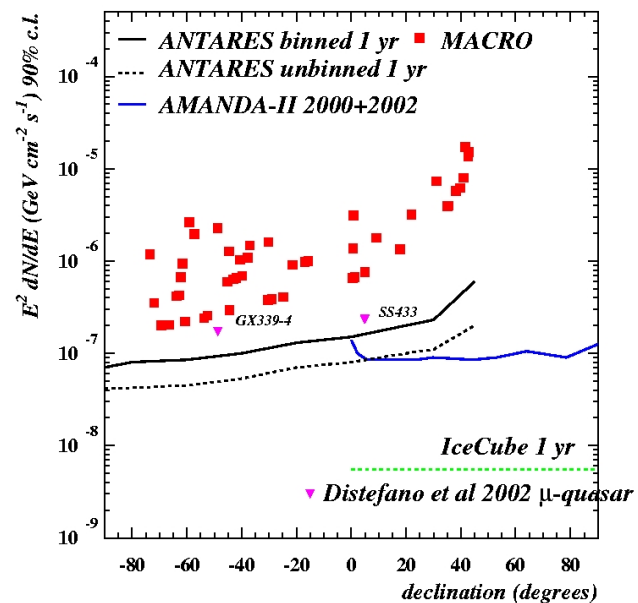
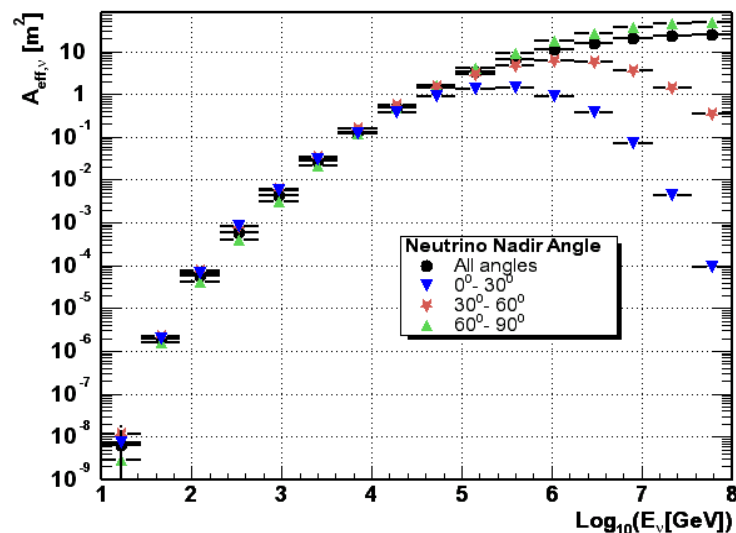
- La **sensibilidad** también se ha calculado en términos del flujo (90% CL):

$$\Phi_{\mu}(\delta = 0) \approx 1.5 \cdot 10^{-15} \text{ cm}^{-2} \text{ s}^{-1} \quad (1 \text{ año})$$

Muon Effective Area



Neutrino Effective Area





Selección del modelo (BIC)

Definición general:

1. Conjunto de datos D
2. Varios modelos $M_1, \dots, M_k$

Teorema de Bayes:

$$p(M_k | D) \propto p(D | M_k) p(M_k)$$

Probabilidad darse M_k dado D =

Probabilidad de M_k de reproducir D
 $\times$
 Probabilidad a priori de M_k

Si $p(M_1) = \dots = p(M_k)$ basta con $p(D|M_k)$ (factor de Bayes)

Integración sobre el espacio de
parámetros desconocidos θ_k

$$p(D | M_k) = \int p(D | \theta_k, M_k) p(\theta_k | M_k) d\theta_k$$

Bayesian Information Criterion o BIC*:

g.d.l (parámetros a ajustar)

$$2\log[p(D | M_k)] \approx 2\log p(D | \hat{\theta}_k, M_k) - v_k \log(n) = \text{BIC}_k$$

$$\text{BIC}_k = 2 \sum_{i=1}^n \log[p_{\text{Gauss}}(\alpha_{RAi}, \delta_i) + p_{\text{BG}}(\delta_i)] - 2 \sum_{i=1}^n \log[p_{\text{BG}}(\delta_i)] - 6 \log(n)$$



Métodos de búsquedas

- Método de bins:

Los métodos de bins buscan una acumulación de sucesos en una pequeña región del cielo (bin). La significancia se estima en comparación con la distribución de fondo.

- Método sin bins:

- 1) Máxima verosimilitud. Maximiza la verosimilitud de una cierta función densidad de acuerdo a dos hipótesis. Una llamada nula (solo fondo) y la hipótesis de una fuente. Se escoge un parámetro (estadística) que es sensible al hecho de que los datos se distribuyan según una hipótesis u otra.