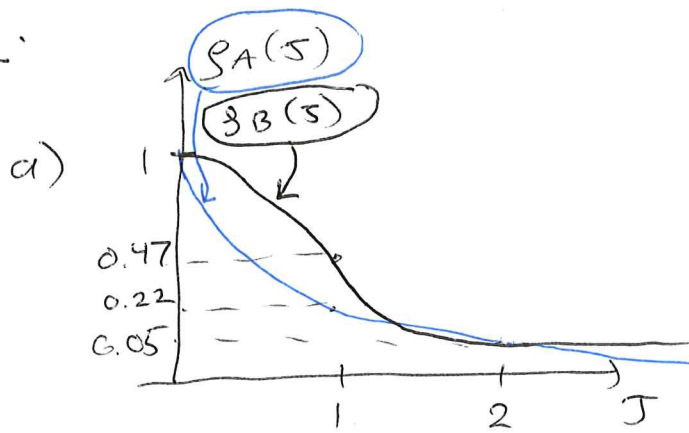


1.



$$\tau=1: e^{-1.5} = 0.22$$

$$- e^{-0.75} = 0.47$$

Effective correlation range $e^{-3} = 0.05$.

$$g^A(\tau) = 0.05 = e^{-1.5 \cdot \tau} \leftrightarrow \tau = \frac{\ln 0.05}{-1.5} = \underline{2}$$

$$g^B(\tau) = 0.05 = e^{-0.75 \tau^2} \leftrightarrow \tau = \sqrt{\frac{\ln 0.05}{-0.75}} = \underline{2}$$

g_A : The exponential correlation function has faster decline of the function near $\tau=0$.

g_B : The Gaussian correlation function is smoother.

The left display is g_B realization. [Left]

The right display is g_A realization. [Right]

1b)

$$E(r(1,1) | r(0,0), r(1,1)) = \hat{r} = \sum_{[(1,1), (0,0)], [(1,1), (1,1)]} \sum_{\substack{(0,0) \\ (1,1)}}^{-1} \begin{pmatrix} d(0,0) \\ d(1,1) \end{pmatrix}$$

PA:

$$\sum_{\substack{(0,0) \\ (1,1)}}^{-1} = \begin{bmatrix} 1 & e^{-1.5\sqrt{2}} \\ e^{-1.5\sqrt{2}} & 1 \end{bmatrix}^{-1} = \begin{pmatrix} \frac{1}{1-e^{-1.5\sqrt{2}}} & -\frac{e^{-1.5\sqrt{2}}}{1-e^{-1.5\sqrt{2}}} \\ -\frac{e^{-1.5\sqrt{2}}}{1-e^{-1.5\sqrt{2}}} & \frac{1}{1-e^{-1.5\sqrt{2}}} \end{pmatrix}$$

$$= \begin{pmatrix} 1.015 & -0.122 \\ -0.122 & 1.015 \end{pmatrix}$$

$$\sum_{[(1,1), (0,0)], [(1,1), (1,1)]} = \begin{pmatrix} e^{-1.5 \cdot 1.56} & e^{-1.5 \cdot 0.14} \end{pmatrix} = (0.097, 0.809)$$

$$\hat{r} = (0.097, 0.809) \cdot \begin{bmatrix} 1.015 \cdot (-2) - 0.122(0.5) \\ -0.122(-2) + 1.015(0.5) \end{bmatrix}$$

$$= \underline{0.405}$$

The exponential correlation has a Markov property, so

$$E(r(1,1,1) | r(1,1)) = e^{-1.5 \cdot 0.14} \cdot \frac{1}{1} \cdot 0.5 = \underline{0.405}$$

It is enough to know $r(1,1)$. $[0.097 \cdot 1.015 - 0.809 \cdot 0.122 = 0]$

ii

1b)

$$S_B: \sum_{\substack{(0,0) \\ (1,1)}}^{-1} = \begin{pmatrix} \frac{1}{1 - e^{-0.75 \cdot 2 \cdot 2}} & - \frac{e^{-0.75 \cdot 2}}{1 - e^{-0.75 \cdot 2 \cdot 2}} \\ - \frac{e^{-0.75 \cdot 2}}{1 - e^{-0.75 \cdot 2 \cdot 2}} & \frac{1}{1 - e^{-0.75 \cdot 2 \cdot 2}} \end{pmatrix}$$

$$= \begin{pmatrix} 1.052 & -0.235 \\ -0.235 & 1.052 \end{pmatrix}$$

$$\sum_{((1,1),(0,0)), ((1,1),(1,1)), (1,1)} = \begin{pmatrix} e^{-0.75 \cdot 2 \cdot 4} & e^{-0.75 \cdot 0 \cdot 0} \end{pmatrix} = (0.163, 0.985)$$

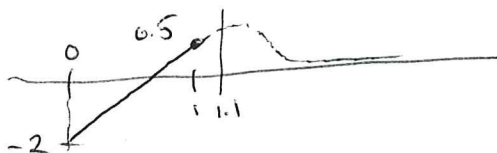
$$\hat{r} = (0.163, 0.985) \begin{bmatrix} 1.052 \cdot (-2) - 0.235(0.5) \\ -0.235(-2) + 1.052(0.5) \end{bmatrix}$$

$$= \underline{0.62}$$

1c)

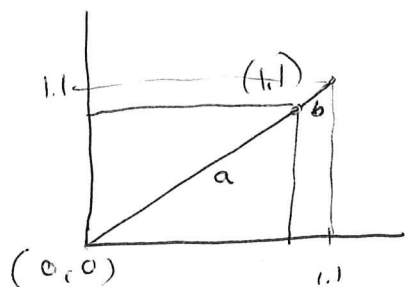
There is no Markov property here.

In fact the prediction is larger than 0.5, because it inherits some of the slope from the -2 measurement at (0,0)



iv)

10)



$$\begin{aligned}
 c &= \|(1,1) - (0,0)\| \\
 &= \|(1,1) - (0,0)\| + \|(1,1) - (1,1)\| \\
 &= a + b
 \end{aligned}$$

$$a = \sqrt{2}, \quad b = \sqrt{2 \cdot 0.1^2} = \sqrt{2} \cdot 0.1$$

$$c = \sqrt{2 \cdot 1.1^2} = \sqrt{2} \cdot 1.1$$

Then

$$\begin{pmatrix} e^{-1.5\sqrt{2} \cdot 1.1} & e^{-1.5\sqrt{2} \cdot 0.1} \end{pmatrix} \begin{pmatrix} 1 \\ -e^{-1.5\sqrt{2}} \end{pmatrix}$$

$$= e^{-1.5\sqrt{2} \cdot 1.1} - e^{-1.5\sqrt{2} \cdot (0.1+1)} = \underline{0}.$$

There is 0 weight on $d(0,0)$
in $\hat{r}$.

The Markov property means that
it is sufficient to know $r(1,1)$.

1d) Prior $\underline{r} = \begin{pmatrix} r(1,1,1) \\ r(0,0) \\ v(1,1) \end{pmatrix} \sim \mathcal{N} \left(\begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}, \underbrace{\begin{pmatrix} 1 & e^{-1.5 \cdot 1.56} & e^{-1.5 \cdot 0.14} \\ e^{-1.5 \cdot 1.56} & 1 & e^{-1.5 \sqrt{2}} \\ e^{-1.5 \cdot 0.14} & e^{-1.5 \sqrt{2}} & 1 \end{pmatrix}}_{\Sigma} \right)$

Likelihood:

$$\underline{d} | \underline{r} \sim \mathcal{N} \left(\underbrace{\begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}}_H \cdot \underline{r}, \begin{pmatrix} 1^2 & 0 \\ 0 & 1^2 \end{pmatrix} \right)$$

Posterior:

$$r | d \sim \mathcal{N} \left(\Sigma H^T [H \Sigma H^T + \Sigma_{d|d}]^{-1} \cdot \underline{d}, \Sigma - \Sigma H^T [H \Sigma H^T + \Sigma_{d|d}]^{-1} H \Sigma \right)$$

We are only interested in $r(1,1,1)$

$$\hat{r} = \begin{pmatrix} e^{-1.5 \cdot 1.56} & e^{-1.5 \cdot 0.14} \end{pmatrix} \begin{bmatrix} 2 & e^{-1.5 \sqrt{2}} \\ e^{-1.5 \sqrt{2}} & 2 \end{bmatrix}^{-1} \cdot \begin{pmatrix} d(0,0) \\ d(1,1) \end{pmatrix}$$

$$= (0.097, 0.809) \cdot \begin{bmatrix} \frac{2}{4 - e^{-1.5 \cdot \sqrt{2} \cdot 2}} & \frac{-e^{-1.5 \sqrt{2}}}{4 - e^{-1.5 \sqrt{2} \cdot 2}} \\ \frac{-e^{-1.5 \sqrt{2}}}{4 - e^{-1.5 \sqrt{2} \cdot 2}} & \frac{2}{4 - e^{-1.5 \sqrt{2} \cdot 2}} \end{bmatrix} \begin{pmatrix} -2 \\ 0.5 \end{pmatrix}$$

$$= (0.097, 0.809) \begin{bmatrix} 0.502 & -0.03 \\ -0.03 & 0.502 \end{bmatrix} \begin{pmatrix} -2 \\ 0.5 \end{pmatrix}$$

$$= \underline{0.153}$$

Compared with 1b) it is closer to 0.
Less influenced by data, more by prior;

1d) Prediction variance:

$$\text{Var}(r(1,1,1.1) | d(0,0), d(1,1))$$

$$= 1 - (0.024, 0.403) \begin{pmatrix} 0.097 \\ 0.809 \end{pmatrix}$$

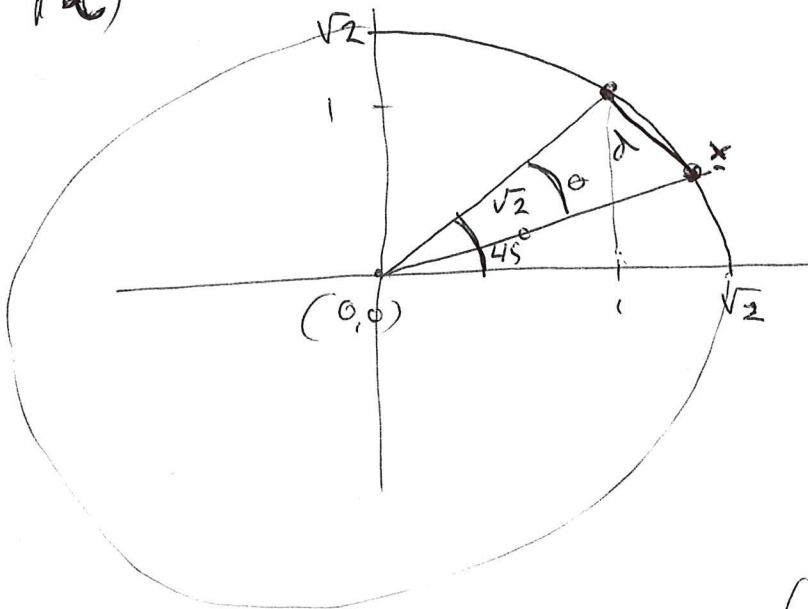
$$= 1 - 0.33 = \underline{0.67}.$$

Formula comes from

$$\text{Var}(\underline{r} | \underline{d}) = \underline{\Sigma}_r - \underline{\Sigma} H^T [H \underline{\Sigma} H^T + \underline{\Sigma}_d]^{-1} H \underline{\Sigma}$$

We are only interested in $r(1,1,1)$ variance.

1e)



Distance from $(0,0)$ to x is $\sqrt{2}$.

Let d be distance from $(1,1)$ to x

$$\hat{r} = \begin{bmatrix} e^{-1.5 \cdot \sqrt{2}} & e^{-1.5 \cdot d} \end{bmatrix} \cdot \begin{pmatrix} 0.502 & -0.03 \\ -0.03 & 0.502 \end{pmatrix} \begin{pmatrix} -2 \\ 0.5 \end{pmatrix}$$

$$= e^{-1.5\sqrt{2}} \cdot (0.502(-2) - 0.03(0.5)) + e^{-1.5d} (-0.03(-2) + 0.502(0.5))$$

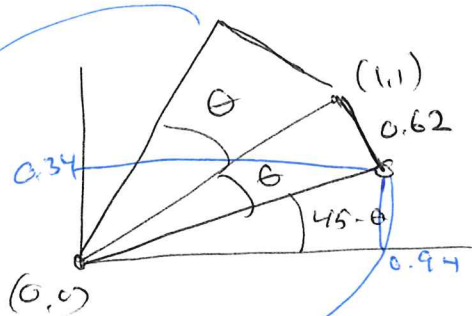
$$= -0.122 + 0.31 \cdot e^{-1.5d}$$

vij

e) For $d=0$, $\hat{r} = -0.122 + 0.31 = \underline{0.189}$.
 For large d , $\hat{r}$ gets negative.
 Solve for $\hat{r} = 0$:

$$-1.5 \cdot d = \ln\left(\frac{0.122}{0.31}\right) = -0.93$$

$$d = \frac{-0.93}{-1.5} = \underline{0.62}$$



The angle is obtained by:

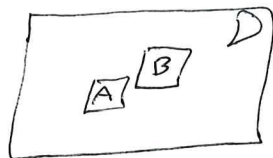
$$0.62 = 2\sqrt{2} \cdot \sin\left(\frac{\theta}{2}\right)$$

$$\theta = 2 \cdot \arcsin\left(\frac{0.62}{2\sqrt{2}}\right) = 0.44 = \underline{25^\circ}$$

$$45 - \theta = \underline{20}$$

$$\underline{x} = \begin{pmatrix} \cos(20) \\ \sin(20) \end{pmatrix} = \begin{pmatrix} 0.94 \\ 0.34 \end{pmatrix}$$

2.
a)



k_A = number of events in $A \subset D$

k_B = number of events in $B \subset D$, $A \cap B = \emptyset$

$$* \begin{cases} k_A \sim \text{Poisson}(\lambda_A) \\ k_B \sim \text{Poisson}(\lambda_B) \end{cases}$$

* k_A and k_B are independent when $A \cap B = \emptyset$

A homogeneous Poisson process has constant intensity parameter λ , such that for a subdomain

$$A \subset D, \quad \lambda_A = \int_{x \in A} \lambda dx = \lambda |A|.$$

Anon-homogeneous Poisson process has spatially varying intensity $\lambda(x)$, such that

$$\lambda_A = \int_{x \in A} \lambda(x) dx.$$

$$b) \quad k_A | k \sim \text{Binomial}\left(k, \frac{\lambda_A}{\lambda_D}\right)$$

$\uparrow$ number of trials $\uparrow$ success probability.

For a homogeneous Poisson process $\frac{\lambda_A}{\lambda_D} = \frac{1}{9}$.

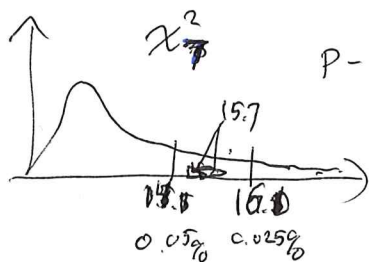
Here, $k = 7 + 9 + 16 + 12 + 12 + 24 + 17 + 11 + 11 = \underline{119}$

2b) The test statistic becomes

$$G = \sum_{i=1}^9 \frac{(O_i - 13.2)^2}{13.2}$$

$$= 2.93 + 1.35 + 0.58 + 0.11 + 0.11 + 8.78 + 1.08 + 0.37 + 0.37$$

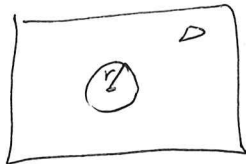
$$= \underline{15.7}$$



p-value ≈ 0.03

This is quite small, but not extremely small (could be random variation).

c)



R_i = distance to nearest event location

$P(R_i > r) = P(0 \text{ events in circle with radius } r)$

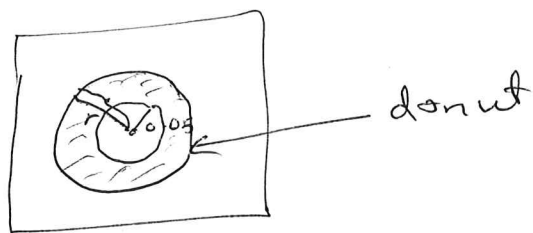
$$\stackrel{\text{Poisson}}{\hat{=}} \frac{(\lambda \pi r^2)^0}{0!} \cdot e^{-\lambda \pi r^2} = e^{-\lambda \pi r^2}$$

Density

$$f_{R_i}(r) = \frac{\partial}{\partial r} [1 - P(R_i > r)] = \underline{2 \lambda \pi r e^{-\lambda \pi r^2}}$$

This is a Weibull distribution with shape parameter 2 and scale parameter $1/\sqrt{2\lambda\pi}$.

c) R_2 = distance to second nearest event location



$$P(R_2 > r | R_1 = 0.05) = P(0 \text{ events in donut})$$

$$= e^{-2\pi(r^2 - 0.05^2)}$$

Density

$$f_{R_2}(r) = \frac{2\lambda\pi r \cdot e^{-2\pi(r^2 - 0.05^2)}}{1} \quad r > 0.05.$$

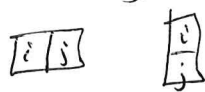
3a)

Joint formulation

$$p(\ell) = \text{const.} \prod_c \psi_c(\ell_j; j \in c)$$

$$= \text{const.} \exp \left\{ \beta \sum_{i,j} I(\ell_i = \ell_j) \right\}$$

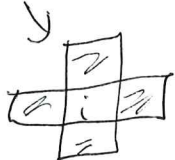
i, j cliques are vertical or horizontal nearest neighbors on grid



Full conditional:

$$p(\ell_i | \ell_j; j \neq i) \underset{\substack{\uparrow \\ \text{Markov}}}{=} p(\ell_i | \ell_j; j \in N_i)$$

N_i : four nearest neighbors



$$\propto \exp \left(\beta \sum_{j \in N_i} I(\ell_i = \ell_j) \right)$$

N_i is neighborhood of cell i .

x)

Algorithms,

A) Gibbs sampling

Initiate

$$\underline{l}^0 = (l_1^0, \dots, l_n^0), \quad n = n_1, n_2$$

Iterate $k = 0, 1, 2, \dots, \text{Iter}$

- draw random cell i

- sample new l_i^* from $p(l_i | l_j^k; j \in N_i)$

$$\text{set } \underline{l}^{k+1} = \underline{l}^k, \quad l_i^{k+1} = l_i^*$$

← using full conditional prob.

B) Metropolis-Hastings

Initiate

$$\underline{l}^0 = (l_1^0, \dots, l_n^0), \quad n = n_1, n_2$$

Iterate $k = 0, 1, 2, \dots, \text{Iter}$

- draw random cell i

$$\text{- set } l_i^* = \begin{cases} 0 & \text{if } l_i^k = 1 \\ 1 & \text{if } l_i^k = 0 \end{cases}$$

- evaluate acceptance rate

$$\alpha = \min \left\{ 1, \frac{p(l_i^* | l_j; j \in N_i) p(l_i)}{p(l_i^k | l_j; j \in N_i) p(l_i)} \right\}$$

$$\text{- } \underline{l}^{k+1} = \underline{l}^k$$

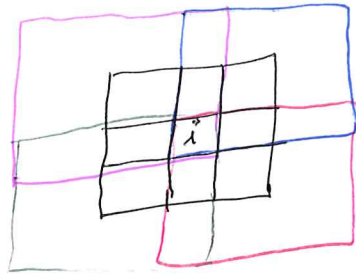
- if $U(0, 1) > \alpha$

$$\text{set } l_i^{k+1} = l_i^*$$

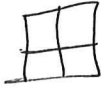
There are other variants of these algorithms.

When $\text{Iter} \rightarrow \infty$, they produce exact (dependent) samples from $p(\underline{l})$.

3b) Second order neighborhood

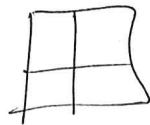


i is in all these neighborhoods,
maximal cliques are

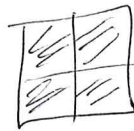
Model $p(l) \propto \exp(\sum_c Q_c(l_c))$  , 2×2 cells.

If symmetry is wanted

2 comb.



all white



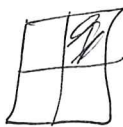
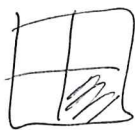
all black

$\leftarrow Q_{\text{equal potential}}$

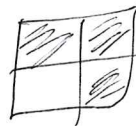
should have same clique potential Q_c

Corners should be the same

8 comb.



$Q_{\text{corner potential}}$



Half should be same vertical/horizontal

4 comb.



$Q_{\text{half potential}}$

Diagonals

2 comb.



$Q_{\text{diag potential}}$

The Gibbs or Metropolis-Hastings sampling is similar (using single-site updating in each iteration).

The difference is in the full conditional probabilities which will now

be $\{q_{\text{equal}}, q_{\text{corner}}, q_{\text{half}}, q_{\text{diag}}\}$

rather than just β or 0 .