



ANSWER BOOK

Hypothesis tests for Poisson and geometric models

Further Statistics 1 · Year FM

6 questions · 20 marks

NAVY EDITION

SECTION A · Warm-up

- A1** $H_0: \lambda = 3$ against $H_1: \lambda > 3$, a one-tailed test. The hypotheses name the population parameter, never the observed count. [2]
- A2** $P(X \geq 8) = 1 - P(X \leq 7) = 0.0119$ assuming H_0 . Since $0.0119 < 0.05$, reject H_0 : there is evidence at the 5% level that the fault rate has risen. [3]

SECTION B · Standard

- B1** $H_0: \lambda = 10$, $H_1: \lambda < 10$, one-tailed. The observed value is low, so use the lower tail: $P(X \leq 4) = 0.0293$. Since $0.0293 < 0.05$, reject H_0 : there is evidence at the 5% level that the accident rate has fallen since the campaign. [4]
- B2** $H_0: p = 0.4$ against $H_1: p < 0.4$. A small p produces a long wait, so the evidence sits in the upper tail of X : $P(X \geq 8) = 0.6^7 = 0.0280$. Since $0.0280 < 0.05$, reject H_0 : there is evidence that the success probability is below 0.4. [4]
- B3** The 5% is split between the two tails, so each carries only 2.5%. The observed tail probability is compared with 0.025 rather than 0.05, which makes rejection harder and can reverse the conclusion of an otherwise identical calculation. [3]

SECTION C · Stretch

- C1** $P(X \geq 8) = 1 - 0.9489 = 0.0511$, which exceeds 0.05, so 8 is not extreme enough. $P(X \geq 9) = 1 - 0.9786 = 0.0214$, which is below 0.05. The critical region is therefore $X \geq 9$, and the test's actual size is 0.0214 rather than the nominal 5%: discreteness always undershoots. [4]