

i will deter state i from violating provided it is moderately effective against state i ; but if the inspection could be used, and would be sufficiently effective, against either state, then it deters them both. The individual states' thresholds are determined by (1.2), and the joint threshold (curved line) by (2.9).

It is not difficult to generalize condition (2.9) to a model in which there are n states to be deterred by the threat of a single inspection. In analogy to (2.5) and (2.6), state i 's unconditional expected payoff is

$$\{[-b_i(1-\beta_i) + d_i\beta_i]q_i\}p_i + d_iq_i(1-p_i) \equiv F_i(q_i, p_i). \quad (2.22)$$

Then it can be shown that, for any $i = 1, 2, \dots, n$, the Nash condition

$$F_i(q_i^*, p_i^*) \geq F_i(q_i, p_i^*) \quad \forall q_i \quad (2.23)$$

is satisfied with $q_i^* = 0$ if and only if

$$\frac{d_i}{b_i + d_i} \frac{1}{1 - \beta_i} \leq p_i^* \quad (2.24)$$

This generalization of (2.10) describes the "Cone of Deterrence" for a single inspection spread over n states. It implies that a necessary and sufficient condition for n states to be deterrable using one inspection is

$$\sum_{i=1}^n \frac{d_i}{b_i + d_i} \frac{1}{1 - \beta_i} \leq 1. \quad (2.25)$$

Condition (2.25) is a further extension of (1.2). Note that the same index,

$$\frac{d_i}{b_i + d_i} \frac{1}{1 - \beta_i}$$

continues to measure "desirability of violating"; as long as the total desirability of violating is low enough and the inspection scheme (the collection of values of p_i) is well-chosen, no violation actually takes place. The threat of inspection is enough to guarantee compliance.

All of the calculations of Problem 2 have been carried out under the assumption that the IAEA must apply all its inspection effort in one state only. In other words, the IAEA selects which state is to be inspected, and then inspects only in that state — it cannot spread some inspection effort over other states. It is appropriate to end with a comment on this assumption.

The models analysed here are interesting because they are natural generalizations of the simple model of Problem 1. Furthermore, the all-or-nothing inspection policy does represent an optimal strategy sometimes (see Problem 3). However there are other cases when it is less than