

Abstract: This paper studies how to design reward mechanisms for cybersecurity-vulnerability verification in web systems. A firm outsources verification tasks to external agents, each of whom holds private information about their own success probability and cost. The firm aims to maximize the probability of discovering a vulnerability while respecting budget and fairness constraints. We characterize the mechanisms that achieve this objective within two classes. The first class consists of mechanisms that satisfy *ex post budget feasibility*, *strategy-proofness*, and *envy-freeness*, whereas the second class also satisfies *interim budget feasibility*. In both classes, the optimal mechanism shares a common structure: it computes each agent's reward-acceptance threshold from the agent's reported type, then uses the thresholds to determine both (i) how many agents are assigned to the verification task and (ii) the corresponding reward amounts.