

Modeling Resources Allocation in Attacker-Defender Games with “Warm Up” CSF

Peiqiu Guan and Jun Zhuang*

Like many other engineering investments, the attacker's and defender's investments may have limited impact without initial capital to “warm up” the systems. This article studies such “warm up” effects on both the attack and defense equilibrium strategies in a sequential-move game model by developing a class of novel and more realistic contest success functions. We first solve a single-target attacker-defender game analytically and provide numerical solutions to a multiple-target case. We compare the results of the models with and without consideration of the investment “warm up” effects, and find that the defender would suffer higher expected damage, and either underestimate the attacker effort or waste defense investment if the defender falsely believes that no investment “warm up” effects exist. We illustrate the model results with real data, and compare the results of the models with and without consideration of the correlation between the “warm up” threshold and the investment effectiveness. Interestingly, we find that the defender is suggested to give up defending all the targets when the attack or the defense “warm up” thresholds are sufficiently high. This article provides new insights and suggestions on policy implications for homeland security resource allocation.

KEY WORDS: Attacker-defender games; contest success functions (CSFs); game theory; subgame-perfect Nash equilibria (SPNE); “warm up” threshold

1. INTRODUCTION

Hundreds of billions of dollars have been spent on homeland security since September 11, 2001,⁽¹⁾ and numerous models^(2–5) have been developed to study the strategic interactions between the governments (defenders) and the terrorists (attackers). In order to help the government to make better decisions in allocating the limited defense resources among multiple targets, binary defense allocation,^(6–12) such as defending or not defending, may not be significantly informative to support the real decision making.

When the defense and the attack efforts are modeled as continuous, instead of binary,

Hirshleifer⁽¹³⁾ and Skaperdas⁽¹⁴⁾ introduce two forms of contest success functions (CSFs) among the players in the field of rent seeking, tournaments, and conflict.⁽¹⁵⁾ One is the ratio form $P(A, D) = \frac{k_1 A^m}{k_1 A^m + k_2 D^m + C}$, and the other is the exponential form $P(A, D) = \frac{\exp[k_1 A]}{\exp[k_1 A] + \exp[k_2 D] + C}$, where $m > 0$ and $k_i > 0$ ($i = 1, 2$) are the mass effect parameters, A and D represent the attacker's and the defender's investment efforts, and C is the inherent defense level.

The CSFs capture the essential relationships among the probability of a successful attack, defense and attack efforts, and the inherent defense levels. The CSFs are normally assumed to be continuous, twice differentiable, and with diminishing marginal returns with respect to both the defense and attack efforts. Table I summarizes the CSFs in the counterterrorism literature. For example, for the function form of CSFs $P(D) = e^{-\lambda D}$, the probability of successful attack $P(D)$ decreases exponentially in

Department of Industrial and Systems Engineering, SUNY at Buffalo, Buffalo, NY, USA.

*Address correspondence to Jun Zhuang, Department of Industrial and Systems Engineering, University at Buffalo, The State University of New York, Buffalo, NY 14260-2050, USA; jzhuang@buffalo.edu.

Table I. CSFs in the Attacker-Defender Game Literature, Where A , D , and C Represent the Attack Effort, Defense Effort, and the Inherent Defense Level, Respectively

Function	References	$\frac{\partial P}{\partial D}$	$\frac{\partial P}{\partial A}$	$\frac{\partial^2 P}{\partial D^2}$	$\frac{\partial^2 P}{\partial A^2}$
Binary attack and continuous defense efforts (Exponential)					
e^{-kD}	Bier <i>et al.</i> , ⁽¹⁶⁾ Hao <i>et al.</i> ⁽¹⁷⁾				
	Wang and Bier ⁽¹⁸⁾	≤ 0	NA	≥ 0	NA
	Shan and Zhuang ⁽¹⁹⁾				
Continuous attack and defense efforts (Ratio)					
$\frac{A}{k(A+D+C)}$	Zhuang and Bier ⁽²⁾	≤ 0	≥ 0	≥ 0	≤ 0
$\frac{A^m}{A^m+D^m}$	Hausken ⁽²⁴⁾	≤ 0	≥ 0	≥ 0	≤ 0
$\frac{A}{A+D+C}$	Hausken and Zhuang ⁽²¹⁾	≤ 0	≥ 0	≥ 0	≤ 0
$\frac{k_1 A}{k_1 A + k_2 D + C}$	Guan <i>et al.</i> ⁽²²⁾	≤ 0	≥ 0	≥ 0	≤ 0
$1 - e^{-kA/D}$	Nikoofal and Zhuang ⁽²³⁾	≤ 0	≥ 0	≥ 0	≤ 0

the defender's effort D ,^(16–19) but it does not depend on the attack effort from the attacker. For another body of the literature, ratio-form CSFs,^(2,20–22) the probability of successful attack decreases convexly in defender's efforts and the inherent defense level, and increases concavely in the attacker's efforts. Nikoofal and Zhuang⁽²³⁾ combine both the ratio and the exponential forms of the CSFs. As shown in Table I, the probability of successful attack decreases in the defender's effort ($\frac{\partial P}{\partial D} \leq 0$) and increases in the attacker's effort ($\frac{\partial P}{\partial A} \geq 0$), both with diminishing marginal returns ($\frac{\partial^2 P}{\partial D^2} \geq 0$, $\frac{\partial^2 P}{\partial A^2} \leq 0$).

Although the property of diminishing marginal returns may hold when the attacker and defender investments are sufficiently high, such property would not hold in practice when the investments are small. For example, depending on specific context, the defender may need to spend millions or billions of dollars to purchase, set up, and test a new security program (e.g., new software to track millions of visitors to the United States). The first several millions of dollars spent may not decrease the probability of a successful attack at all. Similarly, the attacker may have to spend a significant amount of resources to prepare for the attacks, and the initial thousands of dollars (or even millions of dollars in a larger-scale attack plot) may not increase the probability of successful attack.

To our best knowledge, none of the previous literature studies the realistic phenomenon where the diminishing marginal returns over continuous investment levels do not apply. We call such a phenomenon the “warm up” effect. We acknowledge that the term “warm up” could also be interpreted as the time window before the main activity (e.g., security system restarted from a shutdown or sleep mode), and could

be modeled in multiple-period games. By contrast, this article considers “warm up” effects statically, and proposes a new functional form of CSF to model the “warm up” effects in a single-period game, as an extension to the literature provided in Table I.

The rest of the article is organized as follows: Section 2 introduces the notations and assumptions in this article; Section 3 proposes a sequential game model between the attacker and the defender with the “warm up” CSFs, and solves a single-target game model analytically; Section 4 solves for the multiple-target case numerically, illustrates the results with real data, compares the results of the models with and without consideration of the correlation between the level of “warm up” threshold and the investment effectiveness, and provides policy implications for the homeland security resource allocation; Section 5 concludes the article, and the Appendix provides the proofs for propositions.

2. NOTATIONS, ASSUMPTIONS, AND MODELS

2.1. Notations and Assumptions

The notations used throughout the article are defined as follows:

- $n > 0$: The number of targets.
- $G_i \geq 0$: Defense resource allocations to the i th target, $\forall i = 1, 2, \dots, n$.
- $G \equiv (G_1, G_2, \dots, G_n)$: Vector for defense resource allocation.
- $T_i \geq 0$: The attack resources to target i , $\forall i = 1, 2, \dots, n$.

- $T \equiv (T_1, T_2, \dots, T_n)$: Vector for attack resources.
- $A_i > 0$: The inherent defense level of the target i .
- $W_{T_i} \geq 0$ and $W_{G_i} \geq 0$: The “warm up” thresholds for the attack and defense investments for target i , respectively; such threshold is defined as the minimal level of investment before the investment becomes impactful; zero or minimal “warm up” effects could be accounted for by setting the thresholds to be zero.
- $P_i(T_i, G_i) \in [0, 1]$: The probability of a successful attack for target i , which is continuous and decreasing in defensive resource, G_i (when $G_i > W_{G_i}$) with diminishing marginal effect, and increasing in attack resource, T_i (when $T_i > W_{T_i}$) with diminishing marginal effects, $\forall i = 1, 2, \dots, n$:

$$\begin{aligned} \frac{\partial P_i(T_i, G_i)}{\partial G_i} &\leq 0, \quad \frac{\partial^2 P_i(T_i, G_i)}{\partial G_i^2} \geq 0, \\ \frac{\partial P_i(T_i, G_i)}{\partial T_i} &\geq 0, \quad \frac{\partial^2 P_i(T_i, G_i)}{\partial T_i^2} \leq 0. \end{aligned} \quad (1)$$

- k_T and k_G : The effectiveness coefficients of the attacker’s and the defender’s investment “warm up” thresholds, respectively.
- $\beta_i \geq 0$ and $\alpha_i \geq 0$: Effectiveness coefficients of the attack and defense investments to target i , respectively.
- $V_i \geq 0$: Valuation of target i , $\forall i = 1, 2, \dots, n$. For simplicity, we use the same target valuations for both the defender and the attacker⁽⁴⁾ in this article.
- $L_G(T, G)$ and $L_T(T, G)$: The objective functions of the defender and the attacker, respectively.
- $\hat{T}(G) = (\hat{T}_1(G), \hat{T}_2(G), \dots, \hat{T}_n(G))$: Attacker’s best responses in the true model.
- $\tilde{T}(G) = (\tilde{T}_1(G), \tilde{T}_2(G), \dots, \tilde{T}_n(G))$: Attacker’s best responses in the defender’s false belief model.
- (T^*, G^*) : Subgame-perfect Nash equilibria (SPNE) in the true model.
- (T^{**}, G^{**}) : SPNE in the defender’s false belief model.

Note that all the subscripts i will be omitted for the notations in the case of $n = 1$.

Following Azaiez and Bier,⁽²⁵⁾ Wang and Bier,⁽¹⁸⁾ and Shan and Zhuang,⁽²⁶⁾ both the defender and the attacker are assumed to be rational. The in-

teraction between the attacker and the defender is modeled as a sequential game, and the attacker is assumed to be the second mover. The attacker is assumed to choose not to attack if he is indifferent between attacking and not attacking.

2.2. Contest Success Functions

Following Hirshleifer,⁽¹³⁾ Zhuang and Bier,⁽²⁾ Hausken and Zhuang,⁽²⁰⁾ and Hausken and Zhuang,⁽²¹⁾ we consider the ratio-form CSF in this article. Most of the CSFs in the literature assume that the CSFs increase in attack investment and decrease in defense investment, which may not hold in practice when the attack and defense systems need to “warm up.” Different from the literature as summarized in Table I, the “warm up” CSF in this article is defined as a piece-wise ratio function with consideration given to the defense and the attack investment “warm up” effects:

$$\begin{aligned} P_i(T_i, G_i) &= \frac{\beta_i(T_i - W_{T_i})^+}{\beta_i(T_i - W_{T_i})^{++} + \alpha_i(G_i - W_{G_i})^{++} + A_i} \\ &= \begin{cases} 0, & \text{if } T_i \leq W_{T_i} \\ \frac{\beta_i(T_i - W_{T_i})}{\beta_i(T_i - W_{T_i}) + A_i}, & \text{if } T_i > W_{T_i} \text{ \& } G_i \leq W_{G_i} \\ \frac{\beta_i(T_i - W_{T_i})}{\beta_i(T_i - W_{T_i}) + \alpha_i(G_i - W_{G_i}) + A_i}, & \text{if } T_i > W_{T_i} \text{ \& } G_i > W_{G_i} \end{cases} \end{aligned} \quad (2)$$

which has the following properties:

- If G_i is smaller than or equal to the defense “warm up” threshold ($G_i \leq W_{G_i}$), the probability of a successful attack would not be changed by the increase of the defense investment G_i . For example, if the defense investment for an airport screening system is less than the cost of purchasing backscatter machines, the probability of successful attack would remain the same or just slightly decreased in the defense investment.
- If G_i is larger than the defense “warm up” threshold ($G_i > W_{G_i}$), we have $\frac{\partial P_i(T_i, G_i)}{\partial G_i} \leq 0$, which means that the probability of a successful attack decreases in the defense investment.
- If the attack effect T_i is smaller than or equal to the attack “warm up” threshold ($T_i \leq W_{T_i}$), we assume that $P_i(T_i, G_i) = 0$, which means that no attack would be successfully launched. For example, an attacker could not launch a bomb attack successfully if he does not have enough

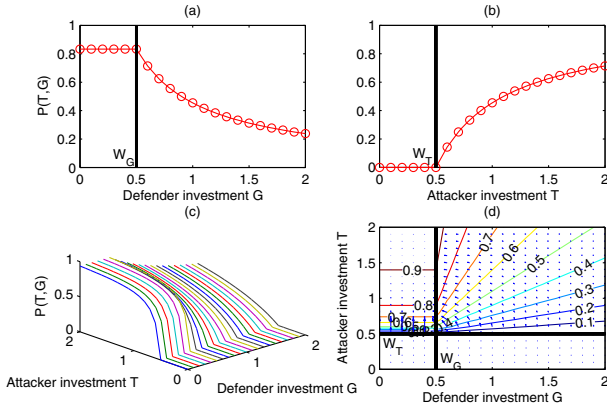


Fig. 1. CSF with consideration of “warm up” effects.

resources to acquire, produce, or use such bomb.

- If T_i is larger than the attack “warm up” threshold ($T_i > W_{T_i}$), we have $\frac{\partial P_i(T_i, G_i)}{\partial T_i} \geq 0$, which means the success attack probability increases in the attack effort.

Fig. 1 illustrates the “warm up” CSF as a function of the attack and defense investments. We only consider a single target in this illustration; i.e., $n = 1$. The baseline values of the model parameters are: $\beta_0 = 1$, $\alpha_0 = 1$, $k_G = k_T = 0$, $A = 0.5$, $G = 1$, $T = 1$, and $W_G = W_T = 0.5$.

Fig. 1(a) shows that the probability of a successful attack decreases in the defender investment G with diminishing marginal effects in the interval of (W_G, ∞) , while the “warm up” CSF remains the same within the defense “warm up” threshold ($G \in [0, W_G]$) as W_G increases. Fig. 1(b) shows that the probability of a successful attack increases in the attacker’s investment T with diminishing marginal effects when $T > W_T$. When the attacker’s resource is less than or equal to the attack “warm up” threshold, $T \in [0, W_T]$, the probability of successful attack equals to zero. Figs. 1(c) and (d) show how the probability of a successful attack changes as both the attack and defense investments vary using a 3-D plot and contour, respectively.

2.3. Modeling Investment Effectiveness Depends on “Warm Up” Threshold

Now we model the scenario in which a higher start-up (“warm up”) cost leads to higher efficiency.⁽²⁷⁾ For example, the backscatter machines for airport screening cost \$250,000 to \$2,000,000

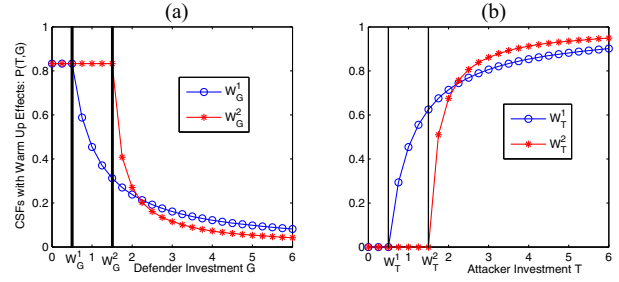


Fig. 2. Probability of a successful attack as functions of defense and attack investments for two levels of defense and attack “warm up” thresholds, $\alpha_0 = \beta_0 = 1$, $k_T = k_G = 1$, $T = G = 1$.

each,⁽²⁸⁾ which is more expensive than the pat-down (no equipment cost), but is less invasive and more efficient. Other examples of such investment include: (1) purchasing of vehicles and unmanned aerial vehicles, and weapons for the border patrollers; (2) purchasing and installing advanced monitors and security systems for federal buildings. This is also true for attack efforts. For example, the dirty bomb attack could cause more damages than the regular bomb attack⁽¹⁵⁾ and it costs the attacker much more “warm up” investment than that in the regular bomb attack. In this article, we model relationships between the investment “warm up” thresholds and their effectiveness coefficients linearly:

$$\beta_i = \beta_0 + k_T W_{T_i}, \quad (3)$$

$$\alpha_i = \alpha_0 + k_G W_{G_i}, \quad (4)$$

where the initial effectiveness coefficients of attack and defense investments are denoted as β_0 and α_0 , and corresponding correlation coefficients are denoted as k_T and k_G ($k_T, k_G \geq 0$), respectively; when $k_T = 0$ or $k_G = 0$, the investment effectiveness coefficients are assumed to be insensitive to the changes of the investment “warm up” threshold; the amounts of the “warm up” thresholds do not impact the efficiency of the attack or the defense systems. When $k_T > 0$, $k_G > 0$, higher investment “warm up” thresholds lead to higher investment effectiveness.

Fig. 2 illustrates the relationships between the “warm up” thresholds and the investment effectiveness. Let W_G^1, W_G^2 ($W_G^1 < W_G^2$) represent two levels of defense “warm up” thresholds; Fig. 2(a) shows that although the probability of a successful attack remains the same for both lines in the intervals of $G \in [0, W_G^1]$ and $G \in [0, W_G^2]$, respectively, the probability of a successful attack decreases at a sharper

rate in the case of W_G^2 than in the case of W_G^1 when $G > W_G^1$ and $G > W_G^2$.

Similarly, Fig. 2(b) shows that the probability of a successful attack remains zero when the attack investment is less than or equal to the attack “warm up” threshold (see the line with when $T \leq W_T^1$, and the line with when $T \leq W_T^2$). The probability of a successful attack increases at a higher rate in the case with attack “warm up” threshold W_T^2 than in the case with attack “warm up” threshold W_T^1 , where $W_T^2 > W_T^1$.

2.4. Optimization Models

In a sequential game model, the attacker is assumed to be a second mover, who can decide the attack strategy T after observing the defender’s resource allocation G over the n targets. The goal of the attacker is to maximize the total expected damage to the defender (CSF weighted by the target valuations), and to minimize the attack costs:

$$L_T(T, G) = \max_T \sum_{i=1}^n \left(\underbrace{\frac{P_i(T_i, G_i)V_i}{\text{Expected damage}}}_{\text{Expected damage}} - \underbrace{T_i}_{\text{Attack costs}} \right). \quad (5)$$

As the first mover, the defender considers the attacker’s best response $\hat{T}(G)$ to the defender’s strategy G , which is defined as

$$\hat{T}(G) \equiv \arg \max_T L_T(T, G), \quad (6)$$

before making the decision. The objective of the defender is to minimize the total expected damage and defense costs:

$$L_G(\hat{T}(G), G) = \min_G \sum_{i=1}^n \left(\underbrace{\frac{P_i(\hat{T}_i(G), G_i)V_i}{\text{Expected damage}}}_{\text{Expected damage}} + \underbrace{G_i}_{\text{Defense costs}} \right). \quad (7)$$

Thus, the SPNE is defined as follows:

Definition 1. We call a collection of strategy (T^*, G^*) an SPNE, if and only if both Equations (8) and (9) are satisfied:

$$T^* = \hat{T}(G^*), \quad (8)$$

$$G^* = \arg \min_G L_G(\hat{T}(G), G). \quad (9)$$

According to the attacker’s best response defined in Equation (6), the SPNE can be solved through backward induction.

3. MODEL SOLUTION AND ANALYSIS

FOR $n = 1$

With the “warm up” CSF defined in Equation (2), this section solves the equilibrium strategies for both the attacker and the defender in the sequential game by using backward induction. We first study the case with a single target $n = 1$.

3.1. Attacker’s Best Response Function

As the second mover in the sequential game model, the attacker’s best response function is given in Proposition 1.

Proposition 1. For $n = 1$, the attacker’s best response function is given by:

$$\hat{T}(G) = \arg \max_T L_T(T, G) = \begin{cases} \text{Case a :} \\ \frac{\sqrt{\beta VA} - A}{\beta} + W_T, & G \leq W_G \& V > \frac{A}{\beta} \& \\ & W_T < \frac{\sqrt{\beta VA} - A}{\sqrt{\beta VA}} V + \frac{A - \sqrt{\beta VA}}{\beta} \\ \text{Case b :} \\ \frac{\sqrt{\beta V \Psi} - \Psi}{\beta} + W_T, & G > W_G \& V > \frac{\Psi}{\beta} \& \\ & W_T < \frac{\sqrt{\beta V \Psi} - \Psi}{\sqrt{\beta V \Psi}} V + \frac{\Psi - \sqrt{\beta V \Psi}}{\beta} \\ \text{Case c :} \\ 0, & G \leq W_G \& V \leq \frac{A}{\beta} \\ & \text{or } G \leq W_G \& V > \frac{A}{\beta} \& \\ & W_T \geq \frac{\sqrt{\beta VA} - A}{\sqrt{\beta VA}} V + \frac{A - \sqrt{\beta VA}}{\beta} \\ \text{Case d :} \\ 0, & G > W_G \& V \leq \frac{\Psi}{\beta} \\ & \text{or } G > W_G \& V > \frac{\Psi}{\beta} \& \\ & W_T \geq \frac{\sqrt{\beta V \Psi} - \Psi}{\sqrt{\beta V \Psi}} V + \frac{\Psi - \sqrt{\beta V \Psi}}{\beta} \end{cases} \quad (10)$$

where $\Psi = \alpha(G - W_G) + A$, $\beta = \beta_0 + k_T W_T$, and $\alpha = \alpha_0 + k_G W_G$.

Remarks: The attacker’s best response function for a single target in Equation (10) shows that the attacker would attack the target in Cases a and b, and put forth zero effort in Cases c and d as his best responses. In Case a, the attacker’s best response effort level does not depend on the value of the defender’s strategy G , while the attacker’s best response effort

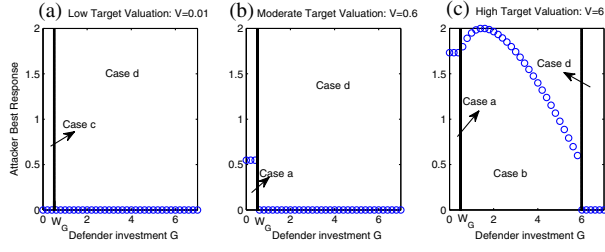


Fig. 3. Best response of the attacker with consideration given to the “warm up” threshold.

first increases then decreases in the defense effort G in Case b.

Cases c and d provide the conditions to deter attacks, including high attack “warm up” threshold, high inherent defense level, high defense investment, or low target valuation.

With the baseline values of model parameters set at $\beta_0 = \alpha_0 = 1$, $k_T = k_G = 0$, $A = 0.5$, $n = 1$, and $W_G = W_T = 0.5$, Fig. 3 illustrates the attacker’s best response function in Proposition 1 with consideration given to both the defender’s and the attacker’s investment “warm up” thresholds, for three levels of target valuations.

Fig. 3(a) shows that the attacker would not attack a low-value target, which illustrates Cases c and d in Proposition 1. Fig. 3(b) shows the attacker attacks a moderate target with the amount of attack resource $\hat{T}(\cdot) = 0.55$ (Case a), which is greater than the attack “warm up” amount $W_T = 0.5$, during the defender’s “warm up” period $G \leq W_G$, and does not attack otherwise (Case d). Fig. 3(c) illustrates the case for the high-value target. The attacker’s best response remains constant (but at a higher level compared to that in Fig. 3(b) for moderate-value target) within the defender’s “warm up” period ($G \leq W_G$) and does not depend on the defense effort G (Case a). When the defense effort G is in a moderate interval ($W_G = 0.5 < G < 6$), the attacker’s best response first increases and then decreases in the defense effort (Case b). We also note that an attack is deterred by a high defense effort G ($G \geq 6$), which is the Case d of the attacker’s best response function in Equation (10).

3.2. Subgame-Perfect Nash Equilibrium (SPNE)

According to Definition 1, substituting the attacker’s best response function in Equation (10) into the defender’s optimization problem (Equation

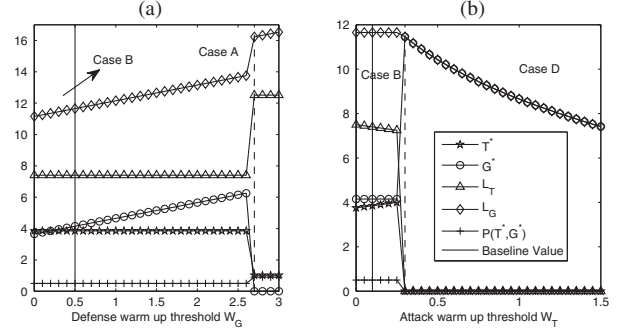


Fig. 4. The sensitivity analyses of W_T and W_G without considering the relationships between the “warm up” thresholds and the investment effectiveness ($k_G = k_T = 0$).

(7)), the SPNE solution is solved and provided in Proposition 2.

Proposition 2. *There are five cases of SPNE solutions (T^*, G^*) for a single-target model. All the SPNE solutions (T^*, G^*) , the corresponding feasible set F^k , optimal conditions O_k , $\forall k = A, B, \dots, E$, the CSF $P(T^*, G^*)$, and attacker’s and defender’s objective functions $L_T(T^*, G^*)$ and $L_G(T^*, G^*)$ are provided in Table II.*

3.3. Sensitivity Analyses

This section studies how the probability of a successful attack, the defender’s and the attacker’s equilibrium strategies, and their objective functions change when the model parameters vary. Based on a set of baseline values, $W_T = 0.1$, $W_G = 0.5$, $\beta_0 = \alpha_0 = 1$, $A = 0.1$, and $V = 15$, we change the same model parameters one at a time and keep the others constant.

We first conduct the sensitivity analyses without consideration of the relationships between the “warm up” thresholds and the investment effectiveness (i.e., $k_G = k_T = 0$).

Fig. 4(a) shows that the defender first increases the defense effort as the defense “warm up” threshold (W_G) increases, and then gives up defending the target if W_G is sufficiently high ($W_G > 2.7$). The defender’s expected damage and costs increase in W_G . From Fig. 4(b), we note that the attack would be deterred by a high attack “warm up” threshold. The defender reduces the defensive resource, even down to zero defense resource, to the target if the attack “warm up” threshold W_T is sufficiently high. The probability of a successful attack slightly decreases as the attack “warm up” threshold increases

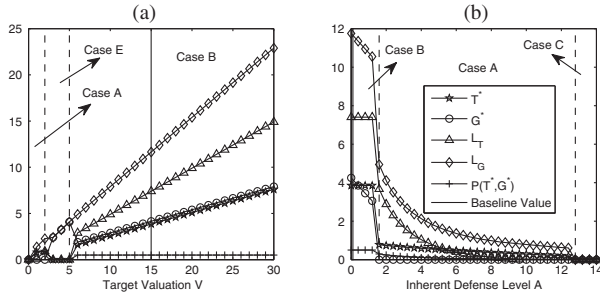


Fig. 5. The sensitivity analysis of V and A without considering the relationships between the “warm up” thresholds and the investment effectiveness ($k_G = k_T = 0$).

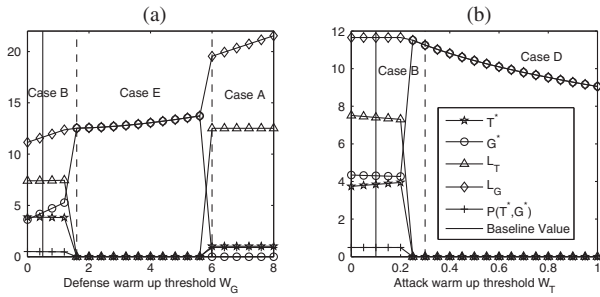


Fig. 6. The sensitivity analysis of W_T and W_G with consideration given to the relationships between the “warm up” thresholds and the investment effectiveness ($k_G = k_T = 0.1$).

and drops to zero when the attack is deterred (for $W_T > 0.3$).

Fig. 5 shows how the model results are sensitive to the change of the target valuation V and the inherent defense level A . In particular, Fig. 5(a) shows that (a) the target with low valuation ($0 < V < 2$) would not be defended but attacked (Case A); (b) the target with moderate valuation ($2 \leq V \leq 5$) would be defended and not be attacked (Case E); and (c) both defender and attacker increase their investment (Case B) if the target’s valuation is large ($V > 5$). From Fig. 5(b), we note that attack could be deterred by a high inherent defense ($A > 12.8$) even when no defense investment is allocated to the target, which implies that no attack would be launched on a target with a high inherent defense level.

Fig. 6 studies the sensitivity analyses of the defense and attack “warm up” thresholds with consideration given to the relationships between the value of the thresholds and the investment effectiveness. In the analyses, we set $k_T = k_G = 0.1$, which means one unit of the “warm up” investments can increase 0.1 units of effectiveness of the attack and defense investments: $\beta = \beta_0 + k_T W_T$ and $\alpha = \alpha_0 + k_G W_G$.

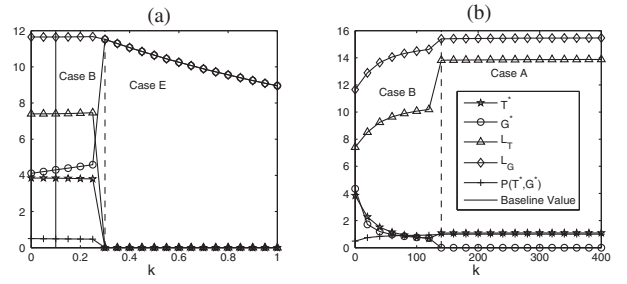


Fig. 7. The sensitivity analyses of k_G, k_T .

Comparing to the case without consideration of the relationships between the defense “warm up” thresholds and the investment effectiveness in Fig. 4, the defender gives up defending with a higher defense “warm up” threshold ($W_G > 6$ in Fig. 6(a), and $W_G > 2.7$ in Fig. 4(a)). Similar to the case without consideration of the relationship between the attack “warm up” threshold and the investment effectiveness in Fig. 4(b), the attack is also finally deterred by some significant level of W_T , which is $W_T > 0.3$.

Fig. 7(a) shows that the defense equilibrium investment G^* first increases and then decreases in k_G (the defense investment becomes more effective). We also note that an attack would be deterred if k_G is high ($k_G > 0.3$) because of the increased effectiveness of the defense investment. Fig. 7(b) shows that both the attack and defense investments (T^* and G^*) decrease as the attack investment becomes more effective, and the defender would give up defending as the k_T is sufficiently high ($k_T \geq 140$).

3.4. Comparison of the Models With and Without Consideration of the Investment “Warm Up” Effects

To study the importance of the novel “warm up” model, we compare the results of the hypothetical model (defender’s false belief model) when the defender does not consider the “warm up” effects ($W_G = W_T = 0$), but in fact the thresholds exist ($W_G > 0, W_T > 0$), with the model proposed in this article (true model). We study the consequence due to this hypothetical belief. In the defender’s belief, her optimal strategy should be defined as follows:

$$\begin{aligned} G^{**} &= \arg \min_G L_G(\bar{T}(G), G) \\ &= P(T, G | W_T = 0, W_G = 0) V + G, \end{aligned}$$

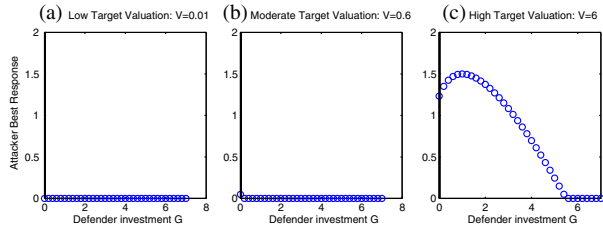


Fig. 8. Best response of the attacker in the hypothetical model where the defender does not consider the “warm up” threshold.

where $\bar{T}(G) = \arg \max L_T(T, G) = P(T, G|W_T = 0, W_G = 0)V - T$ is denoted as the attacker's best response function according to the defender's belief.

However, the attacker makes his best response with consideration given to the investment “warm up” effects. The defender's payoff in the hypothetical mode depends on the defender's equilibrium strategy in the hypothetical model G^{**} and the true attacker equilibrium strategy T^* , which is defined in Equation (8). Thus, due to the false belief, the SPNE (T^{**}, G^{**}) is given as follows:

$$\begin{aligned} T^{**} &= \hat{T}(G^{**}) \\ &= \arg \max L_T(T, G^{**}) = P(T, G^{**})V - T, \quad (11) \\ G^{**} &= \arg \min_G L_G(\bar{T}(G), G) \\ &= P(T, G|W_T = 0, W_G = 0)V + G, \quad (12) \end{aligned}$$

where $\bar{T}(G) = \arg \max L_T(T, G) = P(T, G|W_T = 0, W_G = 0)V - T$. Note that the defender's and the attacker's objective functions at the equilibrium points are denoted as L_G^{**} and L_T^{**} , respectively.

Fig. 8 shows the results of the attacker's best response $\bar{T}(G)$ without considering the “warm up” effects of the attack and defensive investments, $W_T = W_G = 0$.

Comparing the results in Fig. 3, we note that the attacker chooses to not attack a low valued target as his best response, regardless of the defender's investment in both Figs. 3(a) and 8(a). Fig. 8(b) shows that the attacker would only attack a moderate valued target with low attack effort ($T = 0.05$) when the defender's investment is low ($G = 0$), and the attacker would choose not to attack the target if the defender's investment G is large. However, the attacker's attack effort is about 10 times less in Fig. 8(b) as compared to Fig. 3(b) ($T = 0.55$), which implies that the defender would underestimate the attacker's attack effort if she does not consider the “warm up” effects of the investment. Fig. 8(c) shows that the attacker's best response first increases and then de-

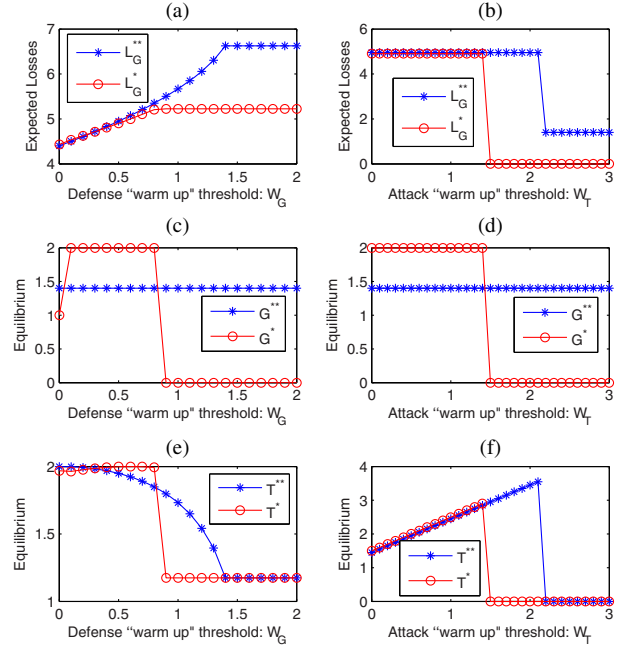


Fig. 9. Comparison of defender's expected payoff, defender's equilibrium strategy, and attacker's equilibrium strategy in the hypothetical model and the true model.

creases as the defender's investment increases for a high valued target case. Comparing the results in Figs. 3(c) and 8(c), we note that the overall attacker's best response would be underestimated if “warm up” effects are not considered. From Fig. 3(c), we also note that the attacker would not attack when his best response level is lower than the attack “warm up” threshold, but in Fig. 8(c) the attacker's best response is positive even when it is below the attack “warm up” threshold, which would lead to the waste of the defense effort.

Figs. 9(a) and (b) show that the defender's payoffs (expected damage and loss) in both models are very close to each other when the defense and attack “warm up” thresholds (W_G and W_T) are small, but the differences enlarge as W_G and W_T increase. In particular, Fig. 9(a) shows that the defender would suffer up to 1.27 (6.62/5.22) times higher expected damage and cost than that in the model proposed in this article (true model) as the defense “warm up” threshold increases if the defender fails to consider the “warm up” effects. In Fig. 9(b), we note that the defender suffers up to 5 units more expected damage and cost than that in the true models if she fails to consider the “warm up” effects.

Figs. 9(c) and (d) show that the defender's equilibrium strategy does not depend on the defense and

attack “warm up” thresholds in the defender’s false belief model, but it does in the true model.

Figs. 9(e) and (f) show that the attacker’s equilibrium strategy is dependent on the defense and attack “warm up” thresholds in both the hypothetical model and the true model. In particular, Fig. 9(e) shows that in the hypothetical model (defender’s false belief model), the attacker’s equilibrium decreases as the defense “warm up” threshold increases when W_G is in a small range ($W_G \leq 1.4$), while in the true model, the attacker’s equilibrium strategy first increases and then decreases in W_G when W_G is in a small range ($W_G \leq 0.9$). The attacker’s equilibrium strategy remains the same in both models when W_G is high (when $W_G > 0.9$ in the true model, and $W_G > 1.4$ in the hypothetical model). Fig. 9(f) shows that the attacker would give up attacking in a higher “warm up” threshold ($W_T = 2.2$) due to the defender’s wrong false belief than that in the true model ($W_T = 1.5$).

4. NUMERICAL ILLUSTRATION FOR MULTIPLE-TARGET MODEL

In this section, we study the equilibrium solution for the model with multiple targets. According to the complex analytical solution for the single-target model in Proposition 2, we expect an intractable solution for the multiple-target case. Instead of obtaining the analytical solution, we focus on the numerical solutions in this section.

4.1. Numerical Illustration

The numerical illustration for the multiple-target model is generated in this section based on the heuristic search algorithm. To illustrate the model with multiple targets, following Bier *et al.*,⁽¹⁶⁾ Hao *et al.*,⁽¹⁷⁾ Hu *et al.*,⁽⁵⁾ Nikoofal and Zhuang,⁽²³⁾ and Shan and Zhuang,⁽¹⁹⁾ we use the expected property damage caused by the terrorist attack for urban areas in the United States to estimate the target valuations.⁽²⁹⁾

In particular, we select the top five most valuable urban areas in the United States, which are New York City, Chicago, San Francisco, Washington, DC, and Los Angeles, as shown in Table III.

For the baseline values of other model parameters, we set $W_G = W_T = 0.5$, $\beta_0 = \alpha_0 = 1$, and $A = 0.1$ for all targets.

We study how the defensive allocation strategies change as the level of defense “warm up” threshold

Table III. Expected Property Damage for the Top Five Urban Areas in the United States

#	Urban Area	Expected Property Losses (V_i \$M)
1	New York City	413.0
2	Chicago	115.0
3	San Francisco	57.0
4	Washington, DC	36.0
5	Los Angeles	34.0

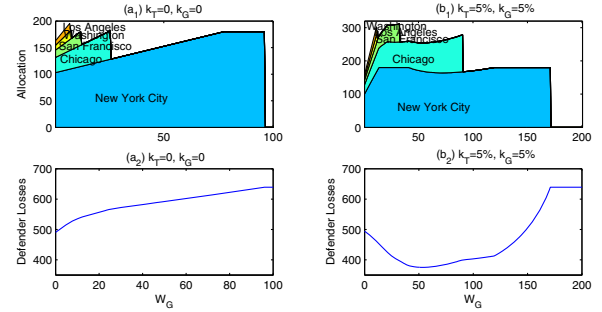


Fig. 10. Defense allocations as functions of the levels of defense “warm up” thresholds levels.

W_G increases, when the relationship of the defense (attack) “warm up” threshold and defense (attack) investment effectiveness coefficients k_G (k_T) are considered in Fig. 10.

Fig. 10 shows that the defender moves the defense resources from the less valuable to more valuable targets as the defense “warm up” threshold increases and finally gives up defending any of the targets if the defense “warm up” thresholds are sufficiently high. Fig. 10 compares optimal defensive resource allocations between two cases: (a) no correlation between “warm up” threshold and investment effectiveness $k_T = k_G = 0$, and (b) low correlation $k_T = k_G = 5\%$. We find that the defender gives up defending targets from the low-valuation targets to the high-valuation target in both cases as the “warm up” threshold increases. For example, the most valuable target “New York City” is the last target that the defender would give up. In Fig. 10(a₂), the defender’s expected damage and cost increase as the defense “warm up” threshold increases when $k_T = k_G = 0$. From Fig. 10(b₂), we note that the defender’s expected damage and cost first decrease and then increase in the case of $k_T > 0$, $k_G > 0$. Since the high defense (attack) “warm up” threshold can induce high defense investment effectiveness in the second case through $\alpha = \alpha_0 + k_G W_G$, as defense investment “warm up” threshold increases, the defender’s investment becomes more effective,

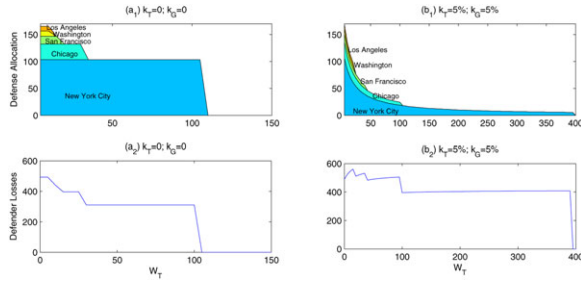


Fig. 11. Defense allocations as functions of the levels of attack “warm up” thresholds levels.

leading to less damage. The expected damage and cost decrease. However, as the defense “warm up” threshold increases, it costs the defender more to “warm up,” leading to a high investment cost. Thus, the defender’s expected damage and costs increase.

Fig. 11 studies how the defender’s defense allocation strategy changes as the attacker’s “warm up” threshold W_T increases. We also compare the optimal attack resource allocations between two cases: (a) no correlation between “warm up” threshold and investment effectiveness $k_T = k_G = 0$, and (b) low correlation $k_T = k_G = 5\%$.

Since a high attack “warm up” threshold can deter the attack, the defender would stop defending the target if the attack is deterred by the attack “warm up” threshold. The defender stops defending targets one by one from the low-valued target to the high-valued target. For example, since the most valuable target, “New York City,” expects the highest attack “warm up” threshold to deter an attack, the defender would stop defending it last. The defender’s expected damage and cost decrease as the attack “warm up” threshold increases for case (a), which is shown in Fig. 11(a₂). In case (b), the defender overall suffers higher expected loss and cost than in case (a) when the correlation between “warm up” threshold and investment effectiveness is considered in Fig. 11(b₂). There are a lot of ways to increase the attack “warm up” threshold, such as increasing the terrorists’ training cost, making it more difficult for them to pass through the security, or acquiring weapons. Though an attack with high attack “warm up” threshold would cause more damage to the defender because of high effectiveness, the attack would be deterred if the attack “warm up” threshold is sufficiently high such that the defender does not need to defend some or all targets. Thus, the defender’s expected damage and cost decrease to zero.

5. CONCLUSIONS AND FUTURE RESEARCH DIRECTIONS

In this article, we study a novel class of contest successful functions (CSF) to capture the “warm up” effects in the attack and defense investments in counterterrorism. To our knowledge, such investment “warm up” effects in attacker-defender resource allocation problems have not been studied in the literature. This article fills the gap by studying attack and defense “warm up” effects in a game-theoretical model.

This article solves the equilibrium solution analytically for the single-target model, identifies five cases of SPNE in the attacker-defender game, and solves the multiple-target model numerically. Interestingly, we find that the defender would give up defending all the targets as either the attack or defense “warm up” thresholds are sufficiently high. For a high defense “warm up” threshold, the defender gives up defending some or all targets because it is too costly to defend the targets. For a high attack “warm up” threshold, the defender stops defending some or all of the targets because the attacks are deterred by the high attack “warm up” threshold, which leads to zero expected damage to the defender. We also find that the defender’s expected damage first decreases and then increases as the defender’s defense “warm up” threshold increases, and it first increases and then decreases as the attacker’s attack “warm up” threshold increases.

This article also provides suggestions on how to allocate limited defense resources to various targets when the defense “warm up” thresholds are considered. We find that not only high inherent defense levels could deter an attack, but also high attack “warm up” thresholds for the attacker. There are scenarios where these are correlated to each other; e.g., for a well-constructed defense system, it usually costs the attacker a high price (attack “warm up” threshold) to pass through the security of the defense system and launch an attack.

In the future, we could consider cooperations among decentralized defenders. For example, if two cities have similar defense needs and are close by, but cannot afford the defense because of the high defense “warm up” threshold, they could share some common defense. Thus, a cooperative defender game among multiple targets with overarching effects could be an interesting extension.

This article considers a complete-information game without deception. In practice, the attacker

may be uncertain about the “warm up” threshold and thus we can study when and how deception and secrecy^(12,30,31) could be used by the defender to mislead the attackers.

In this article, we consider a one-period game where the “warm up” effects are embedded in the CSF, and such effects would not fail. In the future, we could use multi-period or continuous-time games to study more general “warm up” effects, including: (a) a period of heightened vulnerability to the target during the “warm up” period; and (b) the use of backup/standby systems⁽³²⁾ after a potential failure of “warm up.”

ACKNOWLEDGMENTS

This research was partially supported by the U.S. Department of Homeland Security (DHS) through the National Center for Risk and Economic Analysis of Terrorism Events (CREATE) under award number 2010-ST-061-RE0001. This research was also partially supported by the U.S. National Science Foundation under award numbers 1200899 and 1334930. However, any opinions, findings, and conclusions or recommendations in this document are those of the authors and do not necessarily reflect views of the DHS, CREATE, or NSF. We thank the editors and anonymous referees for their helpful comments. The authors assume responsibility for any errors.

APPENDIX

Proof of Proposition 1 A.1

First, we prove the concavity of the attacker’s objective function. For $n = 1$, we have:

$$\max_T L_T(T, G) = P(T, G)V - T.$$

From the assumption in Equation (1), we know that $\frac{\partial^2 P(T, G)}{\partial T^2} \geq 0$. Then we have $\frac{\partial^2 L_T(T, G)}{\partial T^2} = \frac{\partial^2 P(T, G)}{\partial T^2} V \geq 0$.

Thus, the attacker’s maximization problem is concave in T . In order to find a T maximizing the attacker’s objective function, we let the first-order derivative of the attacker’s objective function equal to 0, $\frac{\partial L_T(T, G)}{\partial T} = 0$, and solve for T .

For $n = 1$, according to Equation (2), we rewrite the CSFs as follows:

$$P(T, G) = \begin{cases} 0, & \text{if } T \leq W_T \\ \frac{\beta(T-W_T)}{\beta(T-W_T)+A}, & \text{if } T > W_T \text{ \& } G \leq W_G \\ \frac{\beta(T-W_T)}{\beta(T-W_T)+\alpha(G-W_G)+A}, & \text{if } T > W_T \text{ \& } G > W_G. \end{cases}$$

Then, we rewrite the objective function of the attacker’s optimization problem as follows:

$$\begin{aligned} \max_{T \geq 0} L_T(T, G) \\ = \begin{cases} -T, & \text{if } T \leq W_T \\ \frac{\beta(T-W_T)}{\beta(T-W_T)+A} V - T, & \text{if } T > W_T \text{ \& } G \leq W_G \\ \frac{\beta(T-W_T)}{\beta(T-W_T)+\alpha(G-W_G)+A} V - T, & \text{if } T > W_T \text{ \& } G > W_G. \end{cases} \end{aligned}$$

When $G > W_G$, we have:

$$\begin{aligned} \max_{T \geq 0} L_T(T, G) \\ = \begin{cases} -T, & \text{if } T \leq W_T \\ \frac{\beta(T-W_T)}{\beta(T-W_T)+\alpha(G-W_G)+A} V - T, & \text{if } T > W_T. \end{cases} \quad (A1) \end{aligned}$$

In order to maximize the attacker’s objective when $T \leq W_T$, from Equation (A1), we note that the optimizer is $T^{C_1} = 0$ and the attacker’s objective is $L_T^{C_1} = 0$.

If $T > W_T$, according to the attacker’s objective function defined in Equation (A1), we have $\frac{\partial L_T(T, G)}{\partial T} = 0 \Rightarrow \frac{\beta V [\alpha(G-W_G)+A]}{[\beta(T-W_T)+\alpha(G-W_G)+A]^2} V - 1 = 0$.

Then, we have:

$$\hat{T}(G) = \begin{cases} 0, & \text{if } V \leq \frac{\alpha(G-W_G)+A}{\beta} \\ \frac{1}{\beta} (\sqrt{\beta V [\alpha(G-W_G)+A]}) - [\alpha(G-W_G)+A] + W_T, & \text{if } V > \frac{\alpha(G-W_G)+A}{\beta}. \end{cases}$$

Let $\Psi \equiv \alpha(G-W_G) + A$; we have:

$$\begin{cases} \hat{T}(G)^{C_2} = 0, & \text{if } V \leq \frac{\Psi}{\beta} \\ \hat{T}(G)^{C_3} = \frac{\sqrt{\beta V \Psi} - \Psi}{\beta} + W_T, & \text{if } V > \frac{\Psi}{\beta}. \end{cases} \quad (A2)$$

Substituting Equation (A2) into the attacker’s objective function in Equation (A1), we have:

$$\begin{cases} L_T^{C_2} = 0, & \text{if } V \leq \frac{\Psi}{\beta} \\ L_T^{C_3} = \left(1 - \frac{\Psi}{\sqrt{\beta V \Psi}}\right) V \\ \quad - \left(\frac{1}{\beta} (\sqrt{\beta V \Psi} - \Psi) + W_T\right) & \text{if } V > \frac{\Psi}{\beta}. \end{cases}$$

If all inequality conditions in Equation (A3) hold when $G > W_G$, then $\hat{T}(G)$ is the attacker's best response function:

$$\hat{T}(G) = \begin{cases} 0, & \text{if } V > \frac{\Psi}{\beta} \text{ \& } L_T^{C_1} \geq L_T^{C_3} \\ \hat{T}(G)^{C_3}, & \text{if } V > \frac{\Psi}{\beta} \text{ \& } L_T^{C_3} > L_T^{C_1} \\ 0, & \text{if } V \leq \frac{\Psi}{\beta}. \end{cases} \quad (\text{A3})$$

Thus,

$$\hat{T}(G) = \begin{cases} 0, & \text{if } V > \frac{\Psi}{\beta} \text{ \& } G > W_G \text{ \& } \\ & W_T \geq \frac{\sqrt{\beta V \Psi} - \Psi}{\sqrt{\beta V \Psi}} V + \frac{\Psi - \sqrt{\beta V \Psi}}{\beta} \\ \frac{\sqrt{\beta V \Psi} - \Psi}{\beta} + W_T, & \text{if } V > \frac{\Psi}{\beta} \text{ \& } G > W_G \text{ \& } \\ & W_T < \frac{\sqrt{\beta V \Psi} - \Psi}{\sqrt{\beta V \Psi}} V + \frac{\Psi - \sqrt{\beta V \Psi}}{\beta} \\ 0, & \text{if } V \leq \frac{\Psi}{\beta}. \end{cases} \quad (\text{A4})$$

Similarly, for $G \leq W_G$, according to the CSFs defined in Equation (2), we write the attacker's objective function for $n = 1$ as follows:

$$\begin{aligned} & \max_{T \geq 0} L_T(T, G) \\ & = \begin{cases} -T, & \text{if } T \leq W_T \\ \frac{\beta(T - W_T)}{\beta(T - W_T) + A} V - T, & \text{if } T > W_T. \end{cases} \end{aligned} \quad (\text{A5})$$

From Equation (A5), we note that $T^{C_4} = 0$ is the optimizer for the attacker when $T \leq W_T$, and the corresponding attacker's objective is $L_T^{C_4} = 0$.

If $T > W_T$, according to the attacker's objective function defined in Equation (A5), we have:

$$\begin{aligned} & \frac{\partial L_T(T, G)}{\partial T} = 0 \\ & \Rightarrow \frac{\beta V A}{[\beta(T - W_T) + A]^2} V - 1 = 0 \\ & \Rightarrow \begin{cases} \hat{T}(G)^{C_5} = 0, & \text{if } V \leq \frac{A}{\beta} \\ \hat{T}(G)^{C_6} & \text{if } V > \frac{A}{\beta}. \end{cases} \quad (\text{A6}) \\ & = \frac{1}{\beta} (\sqrt{\beta V A} - A) + W_T, \end{aligned}$$

Substituting the local optimizers in Equation (A6) into the attacker's objective function in Equation

(A5), we have:

$$\begin{cases} L_T^{C_5} = 0, & \text{if } V \leq \frac{A}{\beta} \\ L_T^{C_6} = \left(1 - \frac{A}{\sqrt{\beta V A}}\right) V & \text{if } V > \frac{A}{\beta}. \end{cases} \quad (\text{A7})$$

The attacker's best response function $\hat{T}(G)$ under the condition of $G \leq W_G$ is defined as follows:

$$\begin{aligned} & \hat{T}(G) \\ & = \begin{cases} 0, & \text{if } V > \frac{A}{\beta} \text{ \& } L_T^{C_4} \geq L_T^{C_6} \\ \hat{T}(G)^{C_6}, & \text{if } V > \frac{A}{\beta} \text{ \& } L_T^{C_6} > L_T^{C_4}, \text{ when } G \leq W_G \\ 0, & \text{if } V \leq \frac{A}{\beta} \end{cases} \\ & = \begin{cases} 0, & \text{if } G \leq W_G, V > \frac{A}{\beta} \text{ \& } \\ & W_T \geq \frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V + \frac{A - \sqrt{\beta V A}}{\beta} \\ \frac{\sqrt{\beta V A} - A}{\beta} + W_T, & \text{if } G \leq W_G, V > \frac{A}{\beta} \text{ \& } \\ & W_T < \frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V + \frac{A - \sqrt{\beta V A}}{\beta} \\ 0, & \text{if } G \leq W_G, V \leq \frac{A}{\beta}. \end{cases} \quad (\text{A8}) \end{aligned}$$

Combining Equations (A4) and (A8), the attacker's best response function is summarized as follows:

$$\hat{T}(G) = \arg \max_{T \geq 0} L_T(T, G)$$

$$= \begin{cases} \frac{\sqrt{\beta V A} - A}{\beta} + W_T, & G \leq W_G \text{ \& } V > \frac{A}{\beta} \text{ \& } \\ & W_T < \frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V + \frac{A - \sqrt{\beta V A}}{\beta} \\ \frac{\sqrt{\beta V \Psi} - \Psi}{\beta} + W_T, & G > W_G \text{ \& } V > \frac{\Psi}{\beta} \text{ \& } \\ & W_T < \frac{\sqrt{\beta V \Psi} - \Psi}{\sqrt{\beta V \Psi}} V + \frac{\Psi - \sqrt{\beta V \Psi}}{\beta} \\ 0, & G \leq W_G \text{ \& } V \leq \frac{A}{\beta} \text{ or } \\ & G \leq W_G \text{ \& } V > \frac{A}{\beta} \text{ \& } \\ & W_T \geq \frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V + \frac{A - \sqrt{\beta V A}}{\beta} \\ 0, & G > W_G \text{ \& } V \leq \frac{\Psi}{\beta} \text{ or } \\ & G > W_G \text{ \& } V > \frac{\Psi}{\beta} \text{ \& } \\ & W_T \geq \frac{\sqrt{\beta V \Psi} - \Psi}{\sqrt{\beta V \Psi}} V + \frac{\Psi - \sqrt{\beta V \Psi}}{\beta} \end{cases}$$

where $\Psi \equiv \alpha(G - W_G) + A$, $\beta \equiv \beta_0 + k_T W_T$.

Proof of Proposition 2 A.2

We solve the equilibrium solution for the single-target model by substituting the attacker's best response function (Equation (10)) into the defender's optimization model in Equation (7) under different conditions:

- Case A: Under the conditions of $G \leq W_G$, $V > \frac{A}{\beta}$ & $W_T < \frac{\sqrt{\beta VA} - A}{\sqrt{\beta VA}}V + \frac{A - \sqrt{\beta VA}}{\beta}$, we have the attacker's best response function $\hat{T}(G) = \frac{\sqrt{\beta VA} - A}{\beta} + W_T$. After substituting $\hat{T}(G)$ into the defender's optimization model in Equation (7), we have:

$$\begin{aligned} \min_{G \geq 0} L_G(\hat{T}(G), G) \\ &= \frac{\beta \left(\left(\frac{\sqrt{\beta VA} - A}{\beta} + W_T \right) - W_T \right)}{\beta \left(\left(\frac{\sqrt{\beta VA} - A}{\beta} + W_T \right) - W_T \right) + A} V + G \\ &= \frac{\sqrt{\beta VA} - A}{\sqrt{\beta VA}} V + G. \end{aligned}$$

The minimizer that minimizes $L_G(\hat{T}(G), G) = \frac{\sqrt{\beta VA} - A}{\sqrt{\beta VA}} V + G$ is $G^A = 0$. The feasible set is $F^A \equiv \{V > \frac{A}{\beta} \text{ \& } W_T < \frac{\sqrt{\beta VA} - A}{\sqrt{\beta VA}}V + \frac{A - \sqrt{\beta VA}}{\beta}\}$.

The corresponding attacker's strategy is $T^A = \frac{\sqrt{\beta VA} - A}{\beta} + W_T$. The CSF is $P(T^A, G^A) = \frac{\sqrt{\beta VA} - A}{\sqrt{\beta VA}}$.

And the corresponding payoffs of the attacker and the defender are $L_T^A = \frac{\sqrt{\beta VA} - A}{\sqrt{\beta VA}}V - \frac{\sqrt{\beta VA} - A}{\beta} - W_T$, and $L_G^A = \frac{\sqrt{\beta VA} - A}{\sqrt{\beta VA}}V$.

- Case B: Under conditions of $G > W_G$ & $V > \frac{\Psi}{\beta}$ & $W_T < \frac{\sqrt{\beta V\Psi} - \Psi}{\sqrt{\beta V\Psi}}V + \frac{\Psi - \sqrt{\beta V\Psi}}{\beta}$, we have $\hat{T}(G) = \frac{\sqrt{\beta V\Psi} - \Psi}{\beta} + W_T$. Substitute $\hat{T}(G)$ to the defender's optimization model in Equation (7), we have:

$$\begin{aligned} \min_{G \geq 0} L_G(\hat{T}(G), G) \\ &= \frac{\beta V \left[\left(\frac{\sqrt{\beta V\Psi} - \Psi}{\beta} + W_T \right) - W_T \right]}{\beta \left[\left(\frac{\sqrt{\beta V\Psi} - \Psi}{\beta} + W_T \right) - W_T \right] + \alpha(G - W_G) + A} V + G \\ &= V \left(1 - \sqrt{\frac{\alpha(G - W_G) + A}{\beta V}} \right) + G, \end{aligned} \quad (A9)$$

where $\Psi = \alpha(G - W_G) + A$.

Let $\frac{\partial L_G(\hat{T}(G), G)}{\partial G} = 0$, then we have:

$$\begin{aligned} \frac{\partial L_G(\hat{T}(G), G)}{\partial G} \\ &= V \left(-\sqrt{\frac{1}{\beta V}} \times \frac{\alpha}{2\sqrt{\alpha(G - W_G) + A}} \right) + 1 = 0 \\ \Rightarrow G &= \begin{cases} \frac{\alpha V}{4\beta} - \frac{A}{\alpha} + W_G, & V > \frac{4\beta A}{\alpha^2} \\ 0, & V \leq \frac{4\beta A}{\alpha^2} \end{cases}. \end{aligned}$$

According to the conditions $W_G \geq 0$ and $G > W_G$, we have $G > 0$. Then the local optimizer under this condition is $G^B = \frac{\alpha V}{4\beta} - \frac{A}{\alpha} + W_G$. The corresponding feasible set is $F^B \equiv \{V > \max(\frac{\alpha^2 V}{4\beta^2}, \frac{4\beta A}{\alpha^2}) \text{ \& } W_T < V(1 - \frac{\alpha}{\beta} + \frac{\alpha^2}{4\beta^2})\}$.

The corresponding attacker's strategy is $T^B = V(\frac{\alpha}{2\beta} - \frac{\alpha^2}{4\beta^2}) + W_T$. The CSF is $P(T^B, G^B) = \frac{2\beta V - \alpha}{2\beta V - \alpha + \alpha V}$. And the corresponding attacker's and defender's payoffs are: $L_T^B = (\frac{2\beta V - \alpha}{2\beta V - \alpha + \alpha V} - \frac{\beta\alpha - \alpha^2}{4\beta})V - W_T$, and $L_G^B = (\frac{2\beta V - \alpha}{2\beta V - \alpha + \alpha V} - \frac{\alpha V}{4\beta})V + W_G$.

- Case C: Under the conditions of $G \leq W_G$ & $V > \frac{A}{\beta}$ & $W_T \geq \frac{\sqrt{\beta VA} - A}{\sqrt{\beta VA}}V + \frac{A - \sqrt{\beta VA}}{\beta}$, or $G \leq W_G$ & $V \leq \frac{A}{\beta}$, we have the attacker's best response function $\hat{T}(G) = 0$ such that $\min_{G \geq 0} L_G(\hat{T}(G), G) = G$.

We note that the minimizer $G^C = 0$ minimizes the defender's minimization problem, and the corresponding feasible set is $F^C \equiv \{V > \frac{A}{\beta} \text{ \& } W_T \geq \frac{\sqrt{\beta VA} - A}{\sqrt{\beta VA}}V + \frac{A - \sqrt{\beta VA}}{\beta} \text{ or } V \leq \frac{A}{\beta}\}$. The corresponding attacker's strategy is $T^C = 0$. The CSF is $P(T^C, G^C) = 0$. And the attacker's and defender's corresponding objectives are $L_T^C = 0$, and $L_G^C = 0$.

- Case D: Under the condition of $G > W_G$ & $V > \frac{\Psi}{\beta}$ & $W_T \geq \frac{\sqrt{\beta V\Psi} - \Psi}{\sqrt{\beta V\Psi}}V + \frac{\Psi - \sqrt{\beta V\Psi}}{\beta}$, or $G > W_G$ & $V \leq \frac{\Psi}{\beta}$, we have the attacker's best response function $\hat{T}(G) = 0$ such that $\min_{G \geq 0} L_G(\hat{T}(G), G) = G$.

- Under the condition of $G > W_G$ & $V > \frac{\Psi}{\beta}$ & $W_T \geq \frac{\sqrt{\beta V\Psi} - \Psi}{\sqrt{\beta V\Psi}}V + \frac{\Psi - \sqrt{\beta V\Psi}}{\beta}$,

- * If $V > W_T$, from condition $W_T \geq \frac{\sqrt{\beta V\Psi} - \Psi}{\sqrt{\beta V\Psi}}V + \frac{\Psi - \sqrt{\beta V\Psi}}{\beta}$, we have $\frac{1}{\alpha}(\frac{\beta(V - \sqrt{VW_T})^2}{V} - A) + W_G \leq G \leq \frac{1}{\alpha}(\frac{\beta(V + \sqrt{VW_T})^2}{V} - A) + W_G$. Then the

minimizer is $G^D = \frac{1}{\alpha}(\frac{\beta(V-\sqrt{VW_T})^2}{V} - A) + W_G$.
The corresponding feasible set is:

$$F^D \equiv \left\{ \begin{array}{l} V > W_T, V > \frac{(V-\sqrt{VW_T})^2}{V} + \frac{A}{\beta} \\ W_T \geq 1 - \frac{(V-\sqrt{VW_T})^2}{V} (\frac{V-\sqrt{VW_T}-1}{V-\sqrt{VW_T}}) - V + \sqrt{VW_T} \end{array} \right\}.$$

The corresponding attacker's strategy is $T^D = 0$. The CSF is $P(T^D, G^D) = 0$. And the defender's and defender's corresponding objectives are $L_T^D = 0$, and $L_G^D = \frac{1}{\alpha}(\frac{\beta(V-\sqrt{VW_T})^2}{V} - A) + W_G$.

* If $V \leq W_T$, from condition $W_T \geq \frac{\sqrt{\beta V \Psi} - \Psi}{\sqrt{\beta V \Psi}} V + \frac{\Psi - \sqrt{\beta V \Psi}}{\beta}$, we have $0 \leq G \leq \frac{1}{\alpha}(\frac{\beta(V+\sqrt{VW_T})^2}{V} - A) + W_G$. Then the minimizer is $G = 0$. It conflicts with the conditions $G > W_G, \forall W_G > 0$. So, $G = 0$ is not a feasible solution under this condition.

- Case E: Under the condition of $G > W_G$ & $V \leq \frac{\Psi}{\beta} = \frac{\alpha(G-W_G)+A}{\beta}$, from condition $V \leq \frac{\Psi}{\beta} = \frac{\alpha(G-W_G)+A}{\beta}$, we have $G \geq \frac{\beta V - A}{\alpha} + W_G$. Thus, the minimizer of the case under the condition $G > W_G$ & $V \leq \frac{\Psi}{\beta}$ is $G^E = \frac{\beta V - A}{\alpha} + W_G$. The corresponding feasible set is $F^E \equiv \{V > \frac{A}{\beta}\}$. The corresponding attacker's strategy is $T^E = 0$. The CSF is $P(T^E, G^E) = 0$. And the attacker's and defender's corresponding objectives are $L_T^E = 0$, and $L_G^E = \frac{\beta V - A}{\alpha} + W_G$.

According to the feasible set F^k , case k is optimal if $F^k \cap F^j = \emptyset$, or if $F^k \cap F^j \neq \emptyset$ and $L_G^k \leq L_G^j$, $\forall k, j = A, B, \dots, E, k \neq j$. Therefore, the optimal range of case i is defined as:

$$O_k \equiv \bigcap_{j=A, B, \dots, E, j \neq k} \{ \{F^k \cap F^j \cap \{L_G^k \leq L_G^j\}\} \cup \{F^k \cap \bar{F}^j\} \}, k = A, B, \dots, E.$$

Thus, for Cases A–E, we have the optimal set as follows:

$$O_A = \left\{ \begin{array}{l} \{ \{F^A \cap F^B \cap \{\frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V \leq (\frac{2\beta V - \alpha}{2\beta V - \alpha + \alpha V} - \frac{\alpha V}{4\beta}) V + W_G\} \} \cup \{F^A \cap \bar{F}^B\} \} \\ \{ \{F^A \cap F^C \cap \{\frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V \leq 0\} \} \cup \{F^A \cap \bar{F}^C\} \} \\ \{ \{F^A \cap F^D \cap \{\frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V \leq \frac{1}{\alpha} (\frac{\beta(V-\sqrt{VW_T})^2}{V} - A) + W_G\} \} \cup \{F^A \cap \bar{F}^D\} \} \\ \{ \{F^A \cap F^E \cap \{\frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V \leq \frac{\beta V - A}{\alpha} + W_G\} \} \cup \{F^A \cap \bar{F}^E\} \} \end{array} \right\}$$

$$O_B = \left\{ \begin{array}{l} \{ \{F^B \cap F^A \cap \{(\frac{2\beta V - \alpha}{2\beta V - \alpha + \alpha V} - \frac{\alpha V}{4\beta}) V + W_G \leq \frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V\} \} \cup \{F^B \cap \bar{F}^A\} \} \\ \{ \{F^B \cap F^C \cap \{(\frac{2\beta V - \alpha}{2\beta V - \alpha + \alpha V} - \frac{\alpha V}{4\beta}) V + W_G \leq 0\} \} \cup \{F^B \cap \bar{F}^C\} \} \\ \{ \{F^B \cap F^D \cap \{(\frac{2\beta V - \alpha}{2\beta V - \alpha + \alpha V} - \frac{\alpha V}{4\beta}) V + W_G \leq \frac{1}{\alpha} (\frac{\beta(V-\sqrt{VW_T})^2}{V} - A) + W_G\} \} \cup \{F^B \cap \bar{F}^D\} \} \\ \{ \{F^B \cap F^E \cap \{(\frac{2\beta V - \alpha}{2\beta V - \alpha + \alpha V} - \frac{\alpha V}{4\beta}) V + W_G \leq \frac{\beta V - A}{\alpha} + W_G\} \} \cup \{F^B \cap \bar{F}^E\} \} \end{array} \right\}$$

$$O_C = \left\{ \begin{array}{l} \{ \{F^C \cap F^A \cap \{0 \leq \frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V\} \} \cup \{F^C \cap \bar{F}^A\} \} \\ \{ \{F^C \cap F^B \cap \{0 \leq (\frac{2\beta V - \alpha}{2\beta V - \alpha + \alpha V} - \frac{\alpha V}{4\beta}) V + W_G\} \} \cup \{F^C \cap \bar{F}^B\} \} \\ \{ \{F^C \cap F^D \cap \{0 \leq \frac{1}{\alpha} (\frac{\beta(V-\sqrt{VW_T})^2}{V} - A) + W_G\} \} \cup \{F^C \cap \bar{F}^D\} \} \\ \{ \{F^C \cap F^E \cap \{0 \leq \frac{\beta V - A}{\alpha} + W_G\} \} \cup \{F^C \cap \bar{F}^E\} \} \end{array} \right\}$$

$$O_D = \left\{ \begin{array}{l} \{ \{F^D \cap F^A \cap \{\frac{1}{\alpha} (\frac{\beta(V-\sqrt{VW_T})^2}{V} - A) + W_G \leq \frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V\} \} \cup \{F^D \cap \bar{F}^A\} \} \\ \{ \{F^D \cap F^B \cap \{\frac{1}{\alpha} (\frac{\beta(V-\sqrt{VW_T})^2}{V} - A) + W_G \leq (\frac{2\beta V - \alpha}{2\beta V - \alpha + \alpha V} - \frac{\alpha V}{4\beta}) V + W_G\} \} \cup \{F^D \cap \bar{F}^B\} \} \\ \{ \{F^D \cap F^C \cap \{\frac{1}{\alpha} (\frac{\beta(V-\sqrt{VW_T})^2}{V} - A) + W_G \leq 0\} \} \cup \{F^D \cap \bar{F}^C\} \} \\ \{ \{F^D \cap F^E \cap \{\frac{1}{\alpha} (\frac{\beta(V-\sqrt{VW_T})^2}{V} - A) + W_G \leq \frac{\beta V - A}{\alpha} + W_G\} \} \cup \{F^D \cap \bar{F}^E\} \} \end{array} \right\}$$

$$O_E = \left\{ \begin{array}{l} \{ \{F^E \cap F^A \cap \{\frac{\beta V - A}{\alpha} + W_G \leq \frac{\sqrt{\beta V A} - A}{\sqrt{\beta V A}} V\} \} \cup \{F^E \cap \bar{F}^A\} \} \\ \{ \{F^E \cap F^B \cap \{\frac{\beta V - A}{\alpha} + W_G \leq (\frac{2\beta V - \alpha}{2\beta V - \alpha + \alpha V} - \frac{\alpha V}{4\beta}) V + W_G\} \} \cup \{F^E \cap \bar{F}^B\} \} \\ \{ \{F^E \cap F^C \cap \{\frac{\beta V - A}{\alpha} + W_G \leq 0\} \} \cup \{F^E \cap \bar{F}^C\} \} \\ \{ \{F^E \cap F^D \cap \{\frac{\beta V - A}{\alpha} + W_G \leq \frac{1}{\alpha} (\frac{\beta(V-\sqrt{VW_T})^2}{V} - A) + W_G\} \} \cup \{F^E \cap \bar{F}^D\} \} \end{array} \right\}$$

Thus, the five cases of optimal strategies and their corresponding optimal ranges and payoffs are proved as documented in Table II.

REFERENCES

1. Department of Homeland Security. DHS budget, 2014. Available at: <http://www.dhs.gov/dhs-budget>, Accessed April 1, 2015.
2. Zhuang J, Bier VM. Balancing terrorism and natural disasters-defensive strategy with endogenous attacker effort. *Operations Research*, 2007; 55(5):976–991.
3. Hausken K, Bier VM, Zhuang J. Defending against terrorism, natural disaster, and all hazards. Game theoretic risk analysis of security threats. Pp. 65–97 in *Combining Reliability and Game Theory*, Springer Series on Reliability Engineering, 2008.
4. Golalikhani M, Zhuang J. Modeling arbitrary layers of continuous-level defenses in facing with strategic attackers. *Risk Analysis*, 2011; 31(4):533–547.
5. Hu J, Homem-de Mello T, Mehrotra S. Risk-adjusted budget allocation models with application in homeland security. *IIIE Transactions*, 2011; 43(12):819–839.
6. Lapan HE, Sandler T. To bargain or not to bargain: That is the question. *American Economic Review*, 1998; 78(2): 16–21.
7. Sandler T, Lapan HE. The calculus of dissent: An analysis of terrorists' choice of targets. *Synthese*, 1988; 76(2):245–261.
8. Sandler T, Arce DGM. *Terrorism & game theory*. Simulation & Gaming, 2003; 34(3):319–337.
9. Bier VM, Oliveros S, Samuelson L. Choosing what to protect: Strategic defensive allocation against an unknown attacker. *Journal of Public Economic Theory*, 2007; 9(4):563–587.

10. Dighe NS, Zhuang J, Bier VM. Secrecy in defensive allocations as a strategy for achieving more cost-effective attacker deterrence. *International Journal of Performability Engineering*, 2009; 5(1):31–43.
11. Zhuang J. Impacts of subsidized security on stability and total social costs of equilibrium solutions in an N-player game with errors. *Engineering Economist*, 2010; 55(2):131–149.
12. Zhuang J, Bier VM. Secrecy and deception at equilibrium, with applications to anti-terrorism resource allocation. *Defence and Peace Economics*, 2011; 22(1):43–61.
13. Hirshleifer J. Conflict and rent-seeking success functions: Ratio vs. difference models of relative success. *Public Choice*, 1989; 63(2):101–112.
14. Skaperdas S. Contest success functions. *Economic Theory*, 1996; 7(2):283–290.
15. RMS. Terrorism risk in the post-9/11 era: A 10-year retrospective, 2011. Available at: http://riskinc.com/Publications/9_11_Retrospective.pdf, Accessed April 1, 2015.
16. Bier VM, Haphuriwat N, Menoyo J, Zimmerman R, Culpen AM. Optimal resource allocation for defense of targets based on differing measures of attractiveness. *Risk Analysis*, 2008; 28(3):763–770.
17. Hao M, Zhuang J, Jin S. Robustness of optimal defensive resource allocations in the face of less fully rational attacker. Pp. 886–891 in *Proceedings of the 2009 Industrial Engineering Research Conference*, Miami, FL, 2009.
18. Wang C, Bier VM. Target-hardening decisions based on uncertain multiattribute terrorist utility. *Decision Analysis*, 2011; 8(4):286–302.
19. Shan X, Zhuang J. Cost of equity in homeland security resource allocation in the face of a strategic attacker. *Risk Analysis*, 2013; 33(6):1083–1099.
20. Hausken K, Zhuang J. Defending against a stockpiling terrorist. *Engineering Economist*, 2011; 56(4):321–353.
21. Hausken K, Zhuang J. The timing and deterrence of terrorist attacks due to exogenous dynamics. *Journal of the Operational Research Society*, 2012; 63(6):790–809.
22. Guan P, Zhuang J, Hora SC. Modeling a multi-target attacker-defender game with budget constraints. Working Paper, Department of Industrial and Systems Engineering, University at Buffalo, The State University of New York, 2014.
23. Nikoofal ME, Zhuang J. Robust allocation of a defensive budget considering an attacker's private information. *Risk Analysis*, 2012; 32(5):930–943.
24. Hausken K. Strategic defense and attack for reliability systems. *Reliability Engineering & System Safety*, 2008; 93(11):1740–1750.
25. Azaiez MN, Bier VM. Optimal resource allocation for security in reliability systems. *European Journal of Operational Research*, 2007; 181(2):773–786.
26. Shan X, Zhuang J. Hybrid defensive resource allocations in the face of partially strategic attackers in a sequential defender-attacker game. *European Journal of Operational Research*, 2013; 228(1):262–272.
27. Cobb CW, Douglas PH. A theory of production. *American Economic Review*, 1928; 139–165.
28. Wikipedia. Backscatter X-ray, 2014. Available at: http://en.wikipedia.org/wiki/Backscatter_X-ray, Accessed April 1, 2015.
29. Willis HH, Morral AR, Kelly TK, Medby JJ. Estimating terrorism risk, 2005. Available at: http://www.rand.org/content/dam/rand/pubs/monographs/2005/RAND_MG388.pdf, Accessed April 1, 2015.
30. Zhuang J, Bier VM. Reasons for secrecy and deception in homeland-security resource allocation. *Risk Analysis*, 2010; 30(12):1737–1743.
31. Zhuang J, Bier VM, Alagoz O. Modeling secrecy and deception in a multiple-period attacker-defender signaling game. *European Journal of Operational Research*, 2010; 203(2):409–418.
32. Hausken K. Game theoretic analysis of standby systems. Pp. 77–92 in Holtzman Y (ed). *Advanced Topics in Applied Operations Management*. Rijeka, Croatia: InTech - Open Access Publisher, 2012.