

Avaliação de Confiança Contextual em Grades Computacionais Multimodo usando Plataformas Seguras

Ricardo de Barros Costa, Luiz Fernando Rust da Costa Carmo,

Núcleo de Computação Eletrônica – Universidade Federal do Rio de Janeiro (UFRJ)
Caixa Postal 2324 – 20.010-974 – Rio de Janeiro – RJ – Brasil

ricardocosta@posgrad.nce.ufrj.br, rust@nce.ufrj.br

Abstract. *This paper investigates the joint use of secure platform (TPM), virtualization techniques and trust evaluation techniques in a GRID, for protecting its owns resources and users' applications. Grid is formatted within different operation modes, which one presenting a different protection level in function of its predefined security mechanisms. The operation mode of each application is selected according to the contextual trust exhibited by the user at submission time. Trust calculation is based on four different factors: identification process, platform robustness, communication mode and user's reputation. Simulations results provide subsidies for correlating operation mode and trust thresholds, aiming to optimize the Grid performance.*

Resumo. *Esse artigo investiga o uso integrado de plataformas seguras (TPM), técnicas de virtualização, e mecanismos baseados em confiança, em uma grade computacional, visando à proteção das aplicações, e dos próprios recursos da grade. Para isso, definem-se quatro possíveis modos de operação para o nó de uma grade (multimodo), com níveis de proteção variando em função dos mecanismos de segurança requisitados. O Modo de operação de uma aplicação é selecionado conforme a confiança contextual apresentada pelo usuário no momento de sua submissão. Para o cálculo da confiança são considerados quatro diferentes fatores: forma de identificação, robustez da plataforma cliente, qualidade da comunicação e reputação do usuário. As simulações efetuadas fornecem subsídios quanto aos limiares de níveis de confiança que devem ser usados para a determinação do modo de operação, visando aperfeiçoar o desempenho da grade.*

1. Introdução

Com o crescimento da internet, e da facilidade de compartilhamento de recursos na rede, crescem as potencialidades para o uso de Grades Computacionais (Grid), conseqüentemente, a preocupação quanto aos aspectos de segurança. Torna-se necessário estabelecer um compromisso entre flexibilidade na disponibilização dos recursos e potencialização de riscos de ataques; i.e., quanto menor o controle imposto ao usuário para acesso a estes recursos, maior a probabilidade de acontecer falhas de segurança.

Diversas técnicas têm sido usadas/propostas para prover segurança em sistemas de grade computacional, entre elas: criptografia, controle de acesso, autenticação, PKI,

sandbox, etc. *Sandbox*, por exemplo, é um mecanismo de segurança para execução confiável de programas, normalmente usado em teste de programas provenientes de terceiros não verificados ainda, ou mesmo de usuários não confiáveis. Essencialmente, disponibiliza-se ao usuário um conjunto de recursos controlados rigidamente através de mecanismos como: não persistência de espaço em disco (*scratching*), acesso de rede desabilitado, incapacidade de leitura de dispositivos de entrada, bloqueio das informações do hospedeiro, etc. Em suma, pode-se dizer que o *sandbox* é um tipo específico de virtualização.

O emprego de uma Infra-estrutura de chaves públicas (ICP) é sem dúvidas outro grande aliado na resolução de muitas das demandas relativas à segurança em grade, especificamente no que se refere à autenticação (e.g.: [Globus 2006] define um formato de credencial baseado em X.509). Porém, o uso de ICP limita a flexibilização do uso dos recursos devido às certificações imposta.

Para viabilizar uma maior flexibilidade nas operações com grades computacionais diversas técnicas foram propostas usando processos autômatos que determinam perfis de autorização em função da qualidade da identificação. Essas técnicas avaliam a reputação dos usuários para inferir um nível de confiança, que por sua vez implica em um determinado perfil de autorização. O objetivo é quantificar, o quanto um lado (e.g., a grade) pode confiar nas informações enviadas pelo outro lado (e.g., cliente), quer seja por conhecimento direto (ou seja, que a grade possui do cliente), quer seja pelo conhecimento que outros possuem (informações vindas de outros usuários) – reputação.

Alguns usuários dos recursos da grade podem desejar que suas aplicações não sejam mapeadas em recursos que são mantidos, ou gerenciados, por entidades que eles não confiam. O contrário também é válido, ou seja, os mantenedores dos recursos podem não desejar que estes sejam alocados a usuários que eles não confiem. Vamos supor que o recurso em questão seja processamento. O recurso pode empregar técnicas de *sandbox* para evitar que a tarefa em questão espione (*eavesdropping*) outras atividades computacionais em andamento. De forma similar, a tarefa pode usar criptografia, recolhimento dos dados (*data hiding*), codificação inteligente, ou outro mecanismo para evitar que sua informação sensível seja revelada (*snooping*).

Um fator limitante neste tipo de abordagem é a dependência explícita da qualidade do processo de autenticação, ou seja, ainda é necessário saber com que se está lidando para ser possível deduzir de que ele é capaz. A maioria das propostas prevê o uso de credenciais digitais para uma gerência eficaz do estabelecimento da confiança. O objetivo principal é evoluir o nível de confiança em função de uma troca acumulativa de credenciais. Todavia, isso não evita a possibilidade de fraudes nestas informações.

Por outro lado, com a popularização de co-processadores seguros, como por exemplo, o *Trusted Platform Module* - TPM [TCG 2006], a implementação de funções criptográficas torna-se menos impactante ao usuário, uma vez que o processamento em hardware fornece um alto ganho de performance, e mais segura, devido às próprias características das plataformas seguras, que torna logicamente inacessível o acesso à chave privada.

O uso de plataformas seguras possibilita a geração de informações (resistentes a ataques via software) que podem ser usadas para autenticar um cliente (ou um acesso)

garantindo a autenticidade da parte remota, isto é, garantindo que a mesma não foi comprometida e possui as mesmas características de conexões anteriores. Outro aspecto dessa contribuição é que a validação do cliente passa a envolver tanto o sistema quanto o usuário, pois a mensagem enviada além de ser composta por informações de usuário, também possui informações sobre o sistema operacional (*boot* seguro, aplicação, etc.) que é gerada automaticamente pela plataforma segura.

Esse artigo propõe o uso integrado de plataformas seguras (TPM), técnicas de virtualização, e mecanismos baseados em confiança, em uma grade computacional, visando à proteção das aplicações, e dos próprios recursos da grade. Para isso, define-se uma grade computacional multimodo como sendo composta por nós heterogêneos, apresentando diferentes níveis de proteção em função dos mecanismos de segurança implementados, de acordo com modos de operação previamente alocados. O Modo de operação de uma aplicação (que define um subconjunto de nós destinos elegíveis) é selecionado conforme a confiança contextual apresentada pelo usuário no momento da submissão. Sendo que, a confiança contextual pode ser definida como uma métrica para o risco associado ao contexto do usuário, baseada em quatro diferentes avaliações: forma de identificação, robustez da plataforma do usuário, qualidade da comunicação e reputação do usuário. Em outras palavras, a idéia é preservar a flexibilidade de operação de uma grade computacional, sem nenhuma forma de imposição de mecanismos de segurança a priori, mas fazendo uma avaliação dos riscos inerentes do usuário, e estabelecendo o modo de operação necessário para conviver com estes riscos.

2. Trabalhos Relacionados

[Shen et al. 2006] propõe o sistema *Daonity* (protótipo), baseado em TPM, para suporte à grade através de um *middleware* específico denominado TSS (*Trusted Platform Support Service*). O sistema *Daonity* apenas se preocupa em disponibilizar as funcionalidades intrínsecas do TPM ao sistema em grade.

Abordagens unindo modelos de políticas para controle de acesso e gerenciadores de confiança em Grade são relativamente recentes e visam o estabelecimento da confiança de uma forma gradual e interativa, de forma a atualizar dinamicamente privilégios de acesso. Em [Skogsrud et al. 2003] é proposto um *framework* para negociação da confiança (*Trust-Serv*) para serviço Web. A relação de confiança evolui no tempo via troca de credenciais (cartões de crédito, passaporte e cartões de afinidades) controladas por uma máquina de estados associada à aplicação. As credenciais são avaliadas para habilitar, ou não, a troca de estado e a respectiva alteração do perfil de acesso do usuário. [Tatyana et al. 2005] propõem um *framework* para controle de acesso adaptativo e negociação da confiança que combina uma API de autorização e controle de acesso e um gerenciador de confiança (*Trust Builder*), com o intuito de regular quando, e como, informações sensíveis podem ser reveladas. Esta proposta se caracteriza por uma análise reativa face à ocorrência de falhas, i.e, uma falha (ex. credenciais erradas) implica no aumento do nível de suspeição sobre o usuário, que por sua vez implica em restringir privilégios de acesso.

Os dois trabalhos anteriores fazem uso de credenciais para inferir quanto à confiança no processo de identificação: a diferença básica é que o primeiro é pró-ativo — aumenta a

confiança no sucesso, e o segundo é reativo - aumenta a suspeição (diminui a confiança) no insucesso.

O emprego de um modelo de análise de confiança junto com estratégias de escalonamento de tarefas em Grade é ainda incipiente e o único que se tem conhecimento foi proposto por [Azzedin 2002] com o intuito de evitar que existam pares de usuário/recurso com níveis de confiança incoerentes, por exemplo: que a tarefa de um usuário A, que não confia em um recurso B, não seja alocada exatamente neste recurso B. Este modelo é acionado por três diferentes heurísticas de escalonamento, sendo apresentados alguns resultados comparativos obtidos por simulação. Este trabalho se alinha à proposta deste artigo no sentido de que o valor da confiança influencia no encaminhamento da tarefa de um usuário. Porém, existe uma diferença básica quanto à cláusula condicional usada no escalonamento, uma vez que [Azzedin et al. 2002] usam o nível de confiança relativo - expresso como uma relação entre usuário e recurso, e na nossa proposta o nível de confiança é absoluto (em função de identificação, plataforma, reputação e comunicação); além de que não é previsto uma grade computacional multimodo.

Quanto ao acoplamento do conceito de confiança a uma grade computacional multimodo, a proposta que mais se aproxima é apresentada por [Smith et al. 2006] que elabora uma hierarquia de relacionamentos de confiança em três níveis para interação entre provedores de recursos, produtores de soluções (fornecedor de software/banco de dados) e usuários. Para garantir a confiança nos dois primeiros níveis é proposta uma abordagem do tipo *sandbox*, usando tecnologia de máquina virtual (Xen [Santhanam et al. 2005]) e, para o terceiro nível é proposto o uso de uma solução baseada na plataforma segura do *Trusted Computing Group*. Na verdade, apenas as funcionalidades de atestação remota do TPM são de fato empregadas nesta proposta.

Concluindo, o caráter inovador da abordagem apresentada neste trabalho está apoiado em três fatores principais:

- cálculo da confiança contextual baseada em quatro diferentes avaliações: identificação, plataforma, reputação e comunicação;
- um mecanismo de encaminhamento de tarefas para uma grade computacional multimodo baseado no nível de confiança contextual do usuário;
- especificação dos modos de operação da grade computacional baseado em TPM e técnicas de virtualização.

3. Plataformas Seguras

A mesma facilidade que permite um sistema evoluir o seu software para torná-lo mais eficiente e seguro, permite que um atacante remoto comprometa o mesmo a fim de alterar a sua funcionalidade e esse passe a trabalhar a favor do atacante, quer seja pelo roubo de informações, quer seja pela falsificação das mesmas. Com base na dificuldade de se garantir a confiabilidade e integridade de um sistema (um software) unicamente por meio de um sistema de segurança (outro software), soluções por meio de hardware foram propostas para diminuir e, até mesmo, anular a eficácia de um ataque remoto, trazendo maior segurança aos sistemas existentes.

Entre as diversas iniciativas existentes, destaca-se a do *Trusted Computing Group* (TCG), uma organização sem fins lucrativos criada em 2003, que tem suas origens no

Trusted Computing Platform Alliance (consórcio criado em 1999 entre as empresas Microsoft, IBM, Compaq), pra desenvolver, definir e promover padrões abertos de computação segura baseada em hardware. Inicialmente teve como membros participantes AMD, Hewlett-Packard, IBM, Infineon, Intel, Lenovo, Microsoft e Sun Microsystems. Sua maior contribuição foi o desenvolvimento do *Trusted Platform Module* [TCG 2006], um circuito semicondutor que é capaz de armazenar informação de forma segura, inviolável por ataques via software. O TPM permite o cálculo e armazenamento de *hashes* em registradores internos, conhecidos como *Platform Configuration Registers* (PCR), e facilidades para a geração segura de chaves criptográficas, atestação remota e blindagem de memória.

- *BM* - Blindagem de memória - evita que um código externo ao programa em execução possa ler ou modificar a região de memória reservada para o mesmo;
- *AR* - Atestação Remota - atestação da plataforma feita pela assinatura digital de estruturas internas do TPM pelo uso de AIKs (*attestation identity key*), i.e., de características que são armazenadas nos registros internos do TPM;
- *AS* - Armazenamento Seguro - uma informação ao ser armazenada em disco é criptografada usando criptografia AES e sua chave (criptografada) é armazenada em disco.

Modos de Comunicação previstos:

- *binding* - operação tradicional de criptografia de mensagem pelo uso da chave pública do destino (apenas o destino pode conhecer o conteúdo da mensagem);
- *signing* - associa um espelho da integridade de uma mensagem com a chave usada para gerar sua assinatura (o TPM rotula alguma de suas chaves como “chaves de sinalização” cuja função é criptografar o *hash*);
- *sealing*: é a operação de *binding* condicionada à checagem de métricas da plataforma (especificadas pelo autor da mensagem), de forma que algumas métricas devem ser satisfeitas para que a mensagem possa ser decodificada; a operação de blindagem associa a chave de seção a um conjunto de valores do PCR; uma mensagem blindada é criada selecionando um intervalo de PCRs e criptografando assimetricamente seus valores junto com a chave simétrica usada para codificar a mensagem; de posse da chave privada (criptografia assimétrica), o destinatário apenas recupera a chave simétrica quando o sistema estiver de acordo com os valores PCR solicitados;
- *sealed-signing* – inclui verificação da preservação da integridade na operação de *sealing*.

Apesar de este trabalho focar exclusivamente nas funcionalidades provida pelo TPM, existem outras propostas que complementam esta iniciativa. A Microsoft, com base na proposta do TCG, apresentou o *Next-Generation Secure Computing Base* – NGSCB [Microsoft 2006a], inicialmente conhecida como *Palladium*, que é uma versão Microsoft das regras e implementações do TPM para o desenvolvimento de um sistema operacional seguro. A tecnologia cria um segundo ambiente operacional para proteger o sistema de códigos maliciosos ao fornecer conexões seguras entre aplicativos, periféricos de hardware, memória e armazenamento. O NGSCB inclui um novo componente de software para o Windows denominado *nexus* e um co-processador para operações de criptografia denominado *Security Support Component* (SSC). Com isso, o sistema Operacional fica dividido em dois modos de operação: (i) modo padrão: no qual

o núcleo do Windows e as aplicações existentes executam; e (ii) modo *nexus*: onde o núcleo *nexus* e seus agentes, denominados *Nexus Computing Agents* (NCA) executam em uma região segura de memória. O Modo *nexus* usa memória protegida para garantir: o isolamento do *kernel nexus* dos NCAs, o isolamento dos NCAs entre si e o isolamento dos programas que estão no modo padrão.

A tecnologia LaGrande [Intel 2006] é uma nova tecnologia de segurança nos processadores da Intel (Merom, Conroe e Woodcrest). Essa tecnologia faz uso extensivo do TPM enquanto na manipulação e armazenamento de *hashes*. A tecnologia LaGrande busca criar um camada de proteção para várias operações que costumeiramente seriam desprotegidas. Além dos quatro pilares descritos anteriormente, esta tecnologia prevê ainda: (i) *execução protegida* - permite que programas sejam executados dentro de um ambiente protegido, onde nenhum outro programa pode ter acesso aos recursos utilizados por eles, especialmente a memória RAM – isto é, os dados sendo manipulados e gerados pelo programa (separação de domínio); e (ii) *vídeo protegido* - cria um caminho seguro entre aplicações rodando dentro da área de execução protegida e a memória da placa de vídeo, impedindo que outro programa possa ver ou alterar o que está sendo escrito na tela.

O *Bitlocker Drive Encryption* – BDE [Microsoft 2006b] é um sistema de proteção de dados integrado nas versões *Ultimate* e *Enterprise* do sistema Windows Vista, provendo três modos de operação, dos quais os dois primeiros fazem uso do TPM: (i) Modo de operação transparente - usa o TPM para garantir a blindagem do disco, liberando a carga do sistema somente quando os arquivos de boot estiverem inalterados em relação ao momento geração das chaves; (ii) modo de autenticação de usuário: o usuário deve fornecer algum tipo de autenticação (informação biométrica, etc.) ao ambiente de *pré-boot* a fim de que o carregamento do sistema continue; e (iii) chave USB: a inserção de um dispositivo USB contendo uma chave pré-estabelecida permite a carga do sistema.

4. Avaliação da Confiança Contextual para Grade Multimodo

O conceito de confiança contextual usado neste trabalho pode ser informalmente definido como uma medida de quão certo um provedor de uma grade computacional está a respeito da forma com que o usuário irá se comportar. Para isso, leva-se em consideração informações quanto ao nível de confiança que pode ser atribuída ao processo de identificação, a configuração da plataforma, a reputação e ao processo de comunicação entre o usuário cliente e o gerente da grade computacional.

Para entender a essência da estratégia proposta, usaremos como base a figura 1. Neste cenário o usuário da plataforma *A* deseja submeter à tarefa *T1* à grade. De acordo com o mecanismo de identificação de *T1*, da configuração da plataforma *A*, do modo de comunicação *comA*, e da reputação de *T1*, é calculado um nível de confiança pelo gerente da grade, que é usado para selecionar o modo da grade para o qual a tarefa *T1* é encaminhada. Quanto maior a confiança em *T1*, menor a exigência de controles de segurança no modo de operação da grade.

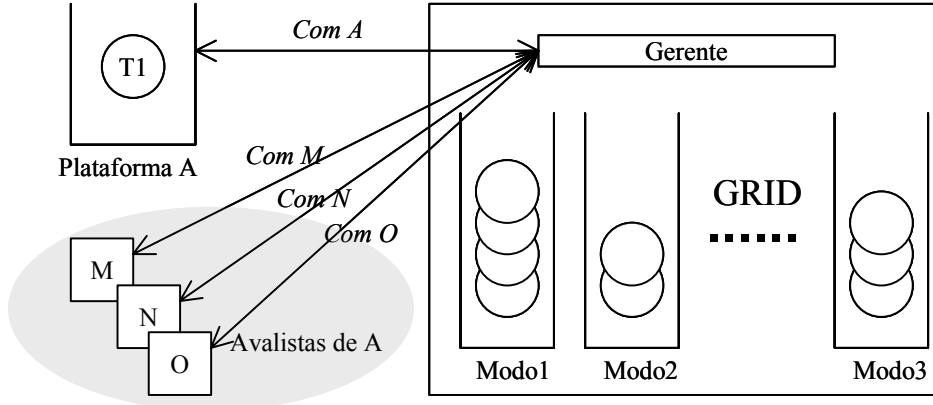


Figura 1: Exemplo de cenário de submissão à grade

Vamos considerar que $T1$ possua certificado digital proveniente de uma autoridade certificadora (CA) reconhecida pela grade, e está usando, portanto uma comunicação assinada. Por outro lado, a plataforma A está configurada em um modo de operação desprotegido, o que implicaria num risco de vazamento da chave privada de A e, da possibilidade de ataques de *personificação* (*masquerade*). É necessário que o nível de confiança atribuído a A capture este risco. Generalizando, podem-se imaginar diferentes situações análogas através da combinação dos fatores de identidade, plataforma, comunicação e reputação. O método proposto neste trabalho quantifica estes aspectos e especifica diferentes modos de operação da grade computacional que devem ser usados em função do nível de confiança resultante. As situações extremas, como falta de confiança e confiança total, correspondem ao modo de operação da grade mais seguro e ao modo com maior desempenho, respectivamente.

Definição 4.1. Um cenário de um Grade é definido pela tripla $(Tarefas^C, Nós^C, \varphi^C)$, onde:

- $Tarefas^C$ é o conjunto de tarefas de C , onde C é uma grade qualquer;
- $Nós^C$ é o conjunto de nós da grade
- φ^C é a *função de atribuição de uma tarefa*, associando cada tarefa a um nó:

$$\varphi^C: Tarefas^C \rightarrow Nós^C$$

Definição 4.2. A distribuição de tarefas entre os nós de uma grade pode ser capturada pelo conjunto denominado:

$$Alocação^C = \{(t, n) \mid \forall t \in Tarefas^C \mid n = \varphi(t)\}$$

Definição 4.3. Sejam *ConfiançaContextual* e *LimiarModoOper* funções com domínios $Tarefas^C$ e $Nós^C$, respectivamente e contradomínios R^+ , uma alocação de tarefas na grade é dita conforme a estratégia de avaliação proposta *se e somente se* a seguinte condição é satisfeita:

$$\forall t \in Alocação^C : ConfiançaContextual(t) > LimiarModoOper(\varphi(t))$$

4.1 Cálculo da Confiança

O valor da confiança contextual de uma tarefa t pode ser expresso pela soma das variáveis CO (comunicação), ID (identificação), PL (configuração da plataforma), e a RE (reputação):

$$ConfiançaContextual(t) = CO^t * ID^t * PL^t * RE^t \quad (1)$$

Após a inicialização, o valor da confiança é atualizado no início e no fim (no caso de incidente) de uma “sessão” entre o usuário e a grade, e toda vez que for indicado como avalista (conforme o procedimento definido em *Reputação* para atualização temporal). No caso de algum incidente durante a execução da tarefa, o valor da confiança contextual é automaticamente zerado.

Identificação - ID

O controle de acesso aos recursos de uma grade computacional tem sido uma das maiores preocupações dos últimos anos. De um lado busca-se uma maior flexibilidade possível para estimular o compartilhamento de recursos, do outro afloram as preocupações quanto à segurança. Buscando garantir maior flexibilidade, os sistemas unindo avaliação de credenciais e reputação tem emergido com muita força. Neste tipo de abordagem, normalmente as credenciais são trocadas de forma a possibilitar certa confiança quanto à identificação do usuário. A confiança, neste tipo de aplicação, substitui integralmente os métodos de autenticação tradicionais, suportando o conceito de usuários livres (sem certificados, logins e senhas). Uma segunda forma de identificação muito empregada (principalmente devido à sua disponibilização em [Globus 2006]) é através de uma Infra-estrutura de chaves públicas (ICP), sendo sem dúvidas um grande aliado na resolução de muitas das demandas relativas à autenticação. Porém, ICP impõe severas limitações na flexibilização do devido às certificações. Por último, existe ainda a possibilidade de se aceitar usuários não-identificados, desde que sejam tomadas severas medidas de segurança.

A questão é que, independente de se impor uma forma de identificação, é importante prever o uso desta diferentes abordagens, atribuindo diferentes níveis de confiança em função da robustez do método. Dessa forma, o método de identificação é expresso pela variável **ID**, que assumirá um valor arbitrário discreto entre 0,1; 0,3 e 0,5; de acordo com um dos três tipos mencionados anteriormente, descritos na tabela 1.

Comunicação - CO

Analogamente ao processo de identificação, avalia-se também a forma com que o usuário troca informações com a grade. A segurança (ou a falta de) no mecanismo de comunicação pode impedir (ou permitir) o sucesso de um ataque. O tipo de comunicação será expresso pela variável **CO**, que pode assumir os seguintes valores discretos: 0,1; 0,3 ou 0,5. Apenas os dois primeiros modos de comunicação previstos pelo TPM são previstos, uma vez que os dois últimos modos de comunicação expressam uma atestação semântica (considerada no contexto de avaliação da plataforma). Além destes dois modos, é inserido um modo desprotegido, para capturar outras situações não previstas.

Tabela 1: Valores atribuídos à variável *ID*

Tipo	Descrição	Peso
Desconhecido	usa alguma informação estruturada desconhecida pela grade, ou não utiliza nenhuma estrutura mais segura de identificação	0,1
Credenciais	garante que a identificação é conhecida da grade, mas a mesma pode ser forjada	0,3
PKI	uso de chaves públicas e certificados digitais reconhecidos por uma CA	0,5

Tabela 2: Valores atribuídos à variável *CO*

Tipo	Descrição	Peso
Desprotegido	sem nenhuma proteção criptográfica	0,1
Binding	operação tradicional de criptografia de mensagem através da chave pública do destino	0,3
Signing	modo Binding + controle de integridade	0,5

Tabela 3: Valores atribuídos à variável *PL*

Tipo	Descrição	Peso
Desprotegido-D	sem nenhuma proteção criptográfica	0,1
Blindagem de memória -B	conjunto de primitivas TPM que garante acesso exclusivo às áreas protegidas (de memória, registrador, etc)	+0,2
Atestação Remota - Ar	atestação semântica da plataforma feita pela assinatura digital dos registros internos do TPM (ex: boot seguro)	+0,2
Armazenamento Seguro -As	uso de criptografia AES para armazenar informações informação em disco	+ 0,2

Plataforma -PL

A métrica plataforma *PL* expressa os controles de segurança implementados pelo usuário, englobando todas as funcionalidades previstas no TPM descritas na seção 3. Note que o uso de atestação remota só pode ocorrer se ambas as partes possuírem o chip TPM e, ainda, combinarem uma atestação semântica a priori. Diferentemente das métricas apresentadas anteriormente, *PL* é expresso através de uma função acumulativa descrita pela equação (2) e representada pela tabela 3. Sejam *B*, *Ar*, *As*, variáveis que assumem o valor 1, quando o respectivo controle definido na tabela 3 for implementado na plataforma de origem da tarefa *t* (e o valor 0, caso contrário), *PL* pode ser definido como:

$$PL^t = 0,1 + +B^t*0,2 + Ar^t*0,2 +As^t*0,2 \quad (2)$$

Reputação - RE

Reputação é a avaliação social que um indivíduo recebe de um conjunto (grupo, sociedade, etc.) e normalmente está associada à confiança que um indivíduo *A* recebe de

um indivíduo B , pelo fato de outros indivíduos com os quais B se relaciona confiar em A . Como mencionado anteriormente, existem diferentes propostas [Resnick et al. 2000], [Guha et al 2004] de como calcular a reputação, sendo que a proposta apresentada a seguir representa uma abordagem sistematicamente acoplada ao conceito de confiança contextual, porém nada impede que sejam usados outros mecanismos de reputação.

O cálculo de reputação é definido pelos seguintes procedimentos:

- o usuário envia à grade uma lista de avalistas (nós em grades ou usuários que já tenham se comunicado antes);
- a grade verifica o nível de confiança atual desses *avalistas* (calculado na última interação de cada um dos avalistas com a grade, ou zerado); fazendo uma atualização temporal (a confiança contextual diminui ou aumenta conforme a frequência de interações entre cliente e grade, modeladas por funções lineares ou exponenciais [Azzedin 2006]);
- sendo L a lista de avalistas, C_l o nível de confiança que a grade possui sobre cada um dos avalista, T_l o intervalo de tempo decorrido a partir da última comunicação com o usuário l , o cálculo da reputação é expresso por:

$$RE^t = \frac{\sum_{\forall L} A(T_l)C_l}{\sum_{\forall L} 1} \quad (3)$$

onde $A(t)$ é uma função de atualização temporal.

O cálculo da confiança contextual não é totalmente imune a ataques do tipo conluio, porém certamente impõe barreiras, uma vez que a “*opinião*” de um usuário é tão mais importante quanto melhor forem os procedimentos de segurança adotados (identificação, comunicação e plataforma) por este.

4.2 Modos de operação da grade

A utilização de uma grade computacional multimodo tem como objetivo principal (e único) a otimização do desempenho dos recursos oferecidos. Partindo da premissa que todo controle de segurança implica numa degradação do desempenho (por menor que seja), é importante que a sua utilização esteja condicionada a um alto risco de exposição a ataques. Essa percepção dos riscos é dada pelo nível de confiança contextual, sendo que, para cada modo de operação previsto nesta seção, deve-se estabelecer um limiar de confiança para sua utilização.

Os modos de operação definidos neste trabalho são definidos pelos modos de operação propostos pelo TPM (considera-se que todos os participantes da grade possuem funcionalidades TPM) em conjunto com técnicas de virtualização (*sandbox*).

Existem diversos grupos de pesquisa em grade que utilizam atualmente *Xen* para prover *sandboxes* seguros nas suas grades computacionais [Keahey et al 2004], [Santhanam et al. 2005], [Barham et al. 2006]. *Xen* fornece máquinas virtuais seguras e independentes, onde instâncias *XenU* rodam geralmente em paralelo numa única máquina física, protegidas uma das outras, sob o controle de um sistema mestre *XenO*. *XenO* é a única instância que tem o acesso ao hardware do sistema, aos sistemas periféricos e aos discos. Os possíveis modos de operação previstos estão definidos na tabela 4.

Tabela 4: Modos de operação de uma grade

Tipo	Descrição	LimiarModoOper
Desprotegido	não utiliza as funções TPM para proteção de dados	L1
Blindado	usa a blindagem de memória – impede leitura e/ou escrita em regiões de memória (registrador, etc.) por código externo	L2
Blindado Seguro	operação no modo blindado em conjunto com o armazenamento seguro de informações (dados criptografados)	L3
SandBox	trabalha no modo blindado seguro com virtualização, para impedir acesso a periféricos, disco	L4

Além da atribuição de recursos baseado no nível de confiança contextual, a grade pode aumentar ou diminuir o modo de segurança (diminuir ou aumentar a performance) total do sistema, solicitando que seus nós alterem o modo de operação, de acordo com uma mudança explícita nos limiares de transição (*L1,L2,L3,L4*).

5. Avaliação da proposta

O provimento de segurança nunca é isento de custos. De acordo com [Suh 2005], o custo para blindagem de memória é, em média, de um crescimento de 6,25% na área de memória ocupada e a perda de desempenho pelas operações de blindagem causam um perda de desempenho da ordem de 5%. As operações de criptografia RSA, ou geração de *hash*, não ocorrem sem perda de desempenho; [Lie et al 2003] apontam uma queda de desempenho total da ordem de 4 a 5% no uso de criptografia da memória por hardware que, somados à operação de blindagem, causa um impacto de 10% na perda de desempenho médio do sistema. Quanto ao uso de *sandbox*, [Menom et al 2005] apresentam um estudo detalhado do desempenho do uso de *Xen* mostrando seu desempenho como sendo da ordem de 80% ao de um sistema do Linux sem o uso de *Xen* e com as mesmas características de hardware. Baseados nestes estudos de desempenho relacionado com os possíveis modos de operação foram estabelecidos os valores contidos na tabela 5.

Tabela 5: Modos de operação x Perda de desempenho

Modo de Operação	Perda Desempenho	Modo de Operação	Perda Desempenho
Desprotegido	0	Blindado Seguro	+10%
Blindado	+5%	Sandbox	+20%

5.1 Metodologia

O principal objetivo da simulação descrita nesta seção é gerar um maior entendimento da relação segurança/desempenho, tendo em vista uma posterior otimização do uso dos modos de operação de uma grade. Apesar do impacto dos limiares dos modos de operação no desempenho da grade ser fortemente influenciado pelo contexto de aplicação, e pelo conjunto de recursos disponíveis, ainda assim é possível estabelecer algumas tendências comportamentais que facilitarão a tarefa de configuração da grade.

O primeiro passo consiste em gerar um conjunto de valores característicos (e significantes) para a variável *confiança contextual*, com o intuito de observar o comportamento dos limiares dos modos de operação em face de diferentes perfis de carga (diferentes usuários na grade simultaneamente). O objetivo é verificar o impacto efetivo no desempenho da grade quando é feito um relaxamento da segurança (limiares maiores) e vice-versa. Para isso, foram geradas todas as combinações (produto) entre as variáveis usadas no cálculo da confiança (*PL*, *CO*, *ID*), sendo selecionadas faixas de valores típicos conforme a tabela 6. Diferentes valores de reputação (0; 0,25; 0,5; 0,75; 1) são aleatoriamente combinados com estes valores.

São retirados os valores que representam combinações improváveis, como por exemplo, um cliente não possuir nenhuma forma de identificação conhecida (ID desconhecido), estabelecer um canal de comunicação desprotegido (CO desprotegido) e possuir mecanismo de atestação remota, que por si só supõe a troca de credenciais e comunicação realizada com mecanismos de proteção. A cada entrada com combinação inválida foi identificada com um “X” na tabela abaixo.

Por fim, os valores da tabela foram normalizados em relação à combinação que apresenta o maior nível de segurança.

Tabela 6: valores típicos de segurança contextual

ID*CO	PL						
	D	B	Ar	As	B+Ar	B+As	B+Ar+As
Desc./ Desp.	0.0067	0.0133	X	X	X	X	X
Desc/Binding	X	0.0400	X	X	X	X	X
Desc./Signing	X	0.0667	X	0.0667	0.1333	X	X
Cred./Desp.	0.0200	0.0400	X	X	X	X	X
Cred./Binding	X	0.1200	0.1200	X	0.2400	X	X
Cred./Signing	X	X	X	0.2000	X	0.4000	0.6000
PKI/Binding	X	X	0.2000	X	0.4000	0.4000	X
PKI/Signing	X	X	X	0.3333	X	0.6667	1.0000

Foram definidos aleatoriamente 500 clientes com modos de operação e níveis de reputação aleatoriamente distribuídos.

Para cada uma das solicitações de serviço de cada cliente foi analisado o risco da grade com base no intervalo de margem de confiança em que o cliente é ser servido e o nível de confiança que o mesmo possui.

Os valores dos limiares de modo de operação (L1, L2, L3, L4) (denominado índice de margem de confiança) foram variados do conjunto de intervalos L1→ [100,15[; L2→ [15,10[; L3→ [10,5[; L4→ [5,0] para o conjunto L1→ [100,90[; L2→ [90,85[; L3→ [85,80[; L4→ [80,0]. De forma a identificar a relação entre o nível de segurança e o nível de desempenho respectivo.

O segundo passo consiste em, a partir desses valores de confiança contextual, estabelecer diferentes configurações de limiares de transição, sendo que para cada configuração é verificado o impacto causado no desempenho.

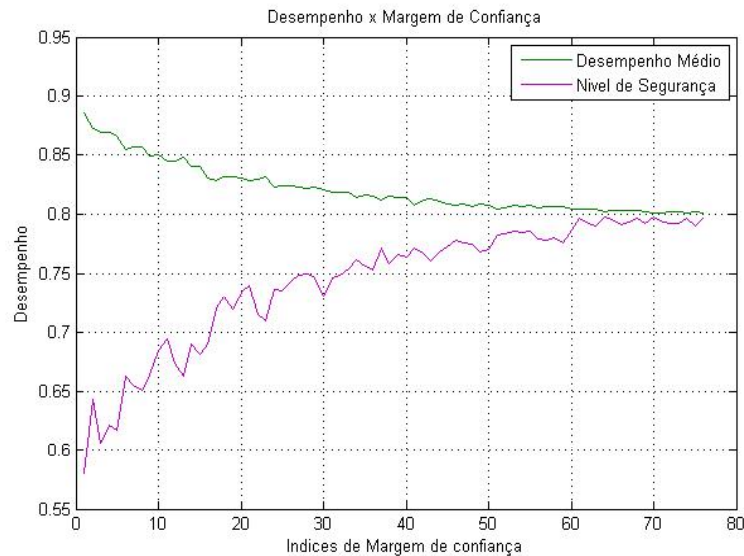


Figura 2: Desempenho médio e Nível de Segurança x Margem de Confiança

5.2 Resultados

O primeiro gráfico da figura 2 apresenta a perda de desempenho médio em relação a uma grade que não esteja utilizando nenhum mecanismo de segurança. O segundo gráfico apresenta uma estimativa do ganho de segurança obtido, permitido que seja possível uma análise comparativa da vantagem efetiva de se implementar tais mecanismos de segurança em grades. Por exemplo, com um aumento no nível de segurança de 65% é possível manter um desempenho próximo de 90%, o que pode ser considerado razoável para muitos sistemas de computação em grade.

6. Conclusões e Trabalhos Futuros

Este artigo apresentou uma proposta para suporte a gestão da segurança em grades computacionais que preserva a flexibilidade de operação, sem imposição do uso de mecanismos específicos a priori por parte do usuário. A estratégia proposta permite uma avaliação dos riscos inerentes a cada usuário pelo cálculo da confiança contextual, e estabelece os requisitos necessários para mitigar estes riscos através da determinação do modo de operação necessário à grade. O cálculo do nível de *confiança contextual* de um usuário é baseado em quatro diferentes fatores: mecanismo de identificação, configuração da plataforma cliente, modo de comunicação e reputação. Este cálculo já considera a possibilidade do uso de plataformas seguras pelo usuário, o que diminui os riscos de ataques e, conseqüentemente, diminuindo as necessidades de controle de segurança na grade.

O artigo propôs ainda diferentes modos de operação para a grade, que por sua vez fazem uso das técnicas mais recentes (e eficientes) para provimento de segurança, i.e., através do emprego de co-processadores seguros e técnicas de virtualização. Os modos

de operação são ativados em função da confiança contextual aferida. A grade computacional multimodo tem como objetivo a otimização do desempenho dos recursos existentes, uma vez que todo controle de segurança implica numa degradação do desempenho. O mecanismo foi avaliado por simulações que fornecem subsídios para a configuração dos limiares de níveis de confiança.

Como trabalho futuro existe a necessidade de um maior aprofundamento na tecnologia LaGrande [Intel 2006], principalmente da proposta do modo de *execução protegida* (prevê um ambiente protegido, onde nenhum outro programa pode ter acesso aos recursos utilizados e aos dados manipulados e gerados por um programa) e a sua potencialidade para gerar funcionalidades do tipo *sandbox* (virtualização). Uma outra frente de trabalho investiga o desenvolvimento de algoritmos para alteração espontânea dos limiares dos modos de operação de uma grade quando confrontado a situações de risco iminente (ataque em curso).

7. Referências

- Azzedin, F., Maheswaran, M. (2002). "Integrating Trust into Grid Resource Management Systems", In the International Conf. on Parallel Processing (ICPP'02).
- Bacon J., Moody K., and Yao, W. (2003). "Access Control and Trust in The Use of Widely Distributed Services", Software-Practice Experience, 33, pp.375–394.
- Barham, P. et al. (2003). "Xen and the Art of Virtualization", In ACM Symposium on Operating Systems Principles (SOSP), USA, pp. 164–177.
- Cheng, A. and Friedman, E. (2005). "Sybilproof reputation mechanisms", In the 2005 ACM SIGCOMM workshop on Economics of peer-to-peer systems - P2PECON.
- Chivers H., Clark J.A. (2004), "Smart dust, friend or foe?—Replacing identity with configuration trust", Computer Networks 46, Elsevier, pp.723–740.
- Foster, I. and Kesselman, C. (2003). "The Grid: Blueprint for a New Computing Infrastructure", Morgan-Kaufmann Publishers Inc., second edition.
- Guha, R., Kumar, R., Raghavan, P., and Tomkins, A., (2004). "Propagation of Trust and Distrust", In the International World Wide Web Conference, New York, USA.
- Intel, (2006). "LaGrande Technology - Preliminary Architecture Specification", disponível em http://www.intel.com/technology/security/downloads/LT_spec_0906.pdf
- Keahey, K., Doering, K., Foster, I. (2004). "From Sandbox to Playground: Dynamic Virtual Environments in the Grid", In the 5th Intern. Workshop on Grid Computing.
- Langheinrich, M, (2003). "When Trust Does Not Compute – The Role of Trust in Ubiquitous Computing", Workshop on Privacy at UBICOMP 2003, Washington.
- Lie, D., Thekkath, C., Horowitz, M., (2003). "Implementing an Untrusted Operating System on Trusted Hardware", In the SOSP'03, New York, USA.
- Lopez, J., Oppliger, R., Pernul, G. (2004). "Authentication and authorization infrastructures (AAIs): a comparative survey", Computers & Security, 23, Elsevier.
- Menon, A., Renato Santos, J., Turner, Y. (2005). "Diagnosing Performance Overheads in the Xen Virtual Machine Environment", In the First ACM/USENIX Conference on Virtual Execution Environments (VEE'05), Chicago, USA.
- Microsoft, (2006a). "Next-Generation Secure Computing Base (NGSCB)", disponível em <http://www.microsoft.com/resources/ngscb/default.mspx>.
- Microsoft (2006b). "Windows BitLocker Drive Encryption Step-by-Step Guide", <http://www.microsoft.com/technet/windowsvista/library/4957-b031-97b4d762cf31.mspx>.

- Platzer C. (2004) "Trust-based Security in Web Services", Master's Thesis, Information Systems Institute, Technical University of Vienna, Austria.
- Resnick, P., Zeckhauser, R., Friedman, E., Kuwabara, K., (2000). "Reputation Systems", Communications of the ACM, Vol 43, Issue 12, pp: 45-48.
- Ryutov, T., Zhou L, Neuman, Foukia, C., Leithead, T., Seamons, K.E. (2005). "Adaptive trust negotiation and access control for GRIDs, In 6th IEEE/ACM International Workshop on Grid Computing – Grid 2005, Washington, USA.
- Santhanam, S., Elango, P., Arpaci-Dusseau, A., Livny, M., (2005). "Deploying Virtual Machines as Sandboxes for the Grid", In the WORLDS '05.
- Shen, Z., Yan, F., Qiang, W., Wu, X., Zhang, H. (2006). Grid System Integrated with Trusted Computing Platform", In the IMSCCS'06.
- Suh, Gookwon E. (2005), "AEGIS: A Single-Chip Secure Processor", Tese de doutorado submetida ao Departamento de Engenharia Elétrica do MIT.
- Skogsrud, H., Benatallah, B., Casati, F. (2003). "Model-Driven Trust Negotiation for Web Services", IEEE Internet Computing, 1089-7801/03, November/December.
- Smith, M., Engel, M., Friese, T., Bernd, F. (2006). "Security Issues in On-Demand Grid and Cluster Computing", In the CCGRIDW'06..
- The Globus Project, (2006). "The Globus Toolkit 4", <http://www.globus.org/toolkit/>.
- Trusted Computing Group – TCG, (2006). "Trusted Platform Module specification", TPM Work Group, disponível em <https://www.trustedcomputinggroup.org/groups/tpm/>.