



A formal method for including the probability of erroneous human task behavior in system analyses

Matthew L. Bolton ^{a,*}, Xi Zheng ^a, Eunsuk Kang ^b

^a University at Buffalo, State University of New York, Department of Industrial and Systems Engineering, Buffalo, NY, USA

^b Institute for Software Research, School of Computer Science, Carnegie Mellon University, NY, USA

ARTICLE INFO

Keywords:

Human error
Formal methods
Model checking
Probabilistic modeling
Human reliability

ABSTRACT

Formal methods have been making inroads into the engineering of human–automation interaction (HAI) by allowing engineers to use mathematical proofs to determine whether normative or unanticipated erroneous human behavior can ever cause problems. However, these approaches are limited because they do not give engineers a way to assess the relative likelihood of different outcomes. In this work, we address this shortcoming by defining a new approach that combines formal approaches with human reliability analysis and probabilistic and statistical model checking. This approach ultimately allows analysts to compute the probability of different outcomes occurring in reactive HAI systems. We describe how this method was realized, assess its scalability, and demonstrate its capabilities with an automated teller machine example. We ultimately discuss our results and describe directions of future research.

1. Introduction

Human error is regularly cited as a source of failure and performance issues in modern systems. It has contributed to more than 1,000,000 injuries and between 44,000 and 98,000 deaths annually in medicine [1]; roughly 75% of all accidents in general aviation and 50% in commercial aviation [2,3]; one third of unmanned aerial system (UAS) accidents [4]; 90% of automobile crashes [5]; and high profile disasters like the accident at Three Mile Island [6]. Humans are often blamed for failures associated with human error. However, the modern perspective holds that errors are the result of shortcomings in system design and thus not solely the fault of human operators, if they are the fault of the human at all. Unfortunately, the complexity of human–automation interaction (HAI) systems and the inherent concurrency that exists in HAI can make it extremely difficult to predict when and how erroneous human behaviors can cause problems. For this reason, a growing body of research has been investigating how rigorous analyses from formal methods can be used to evaluate and design HAI [7–9]. In particular, techniques have been discovered that enable an analyst to pair models of human tasks (a normative description of the behavior humans use to achieve goals when interacting with a system) with models of system functionality. Verification techniques are then used to determine if normative or erroneous human behavior (based on systematic deviations from normative tasks) can result in safety violations [10–15].

These techniques are powerful and well-suited to discovering design flaws in HAIs, especially those involving potentially unanticipated erroneous human behaviors, and thus suggesting interventions that improve system reliability. However, they do have limitations. In particular, they do not account for the relative likelihood of different erroneous behaviors. This situation can make it difficult for analysts to prioritize how to address discovered problems. This research seeks to address this deficiency by extending formal, task-based, verification methods with concepts from human reliability analysis (HRA) and probabilistic/statistical model checking (methods for formally verifying stochastic systems). In what follows, we cover the background necessary for understanding our approach, we describe our method, we analyze a proof of concept example [an automated teller machine (ATM)] to check that our method is providing valid results. Finally, we discuss our developments and discuss directions for future work.

2. Background

2.1. Formal methods

Formal methods is a computer science sub-domain concerned with mathematically modeling systems with precise languages, specifying desirable system properties, and verifying (proving) whether the properties hold with the system. Model checking [16] is computer software

* Corresponding author.

E-mail address: mbolton@buffalo.edu (M.L. Bolton).

that performs formal verification automatically, where efficient data structures and algorithms are used to search through a system model's state space to exhaustively determine if specification properties hold.

The majority of formal methods are non-stochastic. Emerging techniques such as probabilistic and statistical model checking enable analysts to account for probabilities in formal verification [17]. In these, stochastic models (such as Markov chains) are used to describe system behavior and the specification properties either describe desirable system properties (which can include probabilities) or request that a probability be computed. Verification then proves whether the specification property holds for the entire model or computes the probability requested in the property. Statistical model checking only differs from probabilistic checking in that, instead of verifying a property against the entire system model, statistical checking checks the property against a number of samples/traces through the model. This produces approximate results. For situations where the property is computing a probability, this means that statistical checking produces a probability value and an accurate confidence interval (CI) for it.

2.2. Formal and task-analytic methods

Formal methods (and especially model checking) are adept at discovering unexpected interactions that can cause system failures. Because of this, a growing literature has been investigating how formal methods can be used to engineer HAI. Comprehensive reviews can be found in [7–9]. Due to topicality, what follows focuses on formal methods that have used task analytic methods and as well as those that have attempted to model human error stochastically.

Task analysis is a systematic process human factors engineers use to describe how humans normatively achieve goals with a system as a task model [18]. Task models can be interpreted formally and thus included in larger formal analyses. This means that formal verification can prove whether human behavior captured in task models can cause performance and safety issues and be used to investigate solutions to problems. Task models can be formally modeled manually [19,20]. However, it is more common to automatically translate tasks represented in their own notations into a formalism where other system behavior are described [10,11,14,21–23].

Additionally, researchers have discovered a number of approaches for injecting or generating human error into previously normative task models so that the impact of both anticipated and potentially unanticipated erroneous behaviors can be accounted for in verifications. These approaches either use mutation patterns (which are manually applied to task models by analysts) or different theories or taxonomies of human error to automatically include deviations from task in formal representations [13–15,24–30].

Because of its support for including human behavior in formal verification analyses, especially as it relates to the automated generation of potentially unanticipated human errors, we use the enhanced operator function model (EOFM) for task modeling in this research. We discuss EOFM next.

2.2.1. The enhanced operator function model

EOFM [23] is an XML-based task modeling formalism where human behavior is captured as an input/output model. Inputs are system elements external to the human (e.g., interface display information and observable environment elements). Outputs are human actions. The operators' tasks describe how human actions occur based on input and local variables, where local variables represent the human's perceptual or cognitive state.

Tasks are a hierarchy of activities and actions (acts) starting with a root activity. Any activity (which represents a goal-directed complex behavior) can decompose into sub-activities and, ultimately, atomic actions (concrete cognitive or observable behaviors that are not further broken down). Strategic knowledge (Boolean expressions) can be specified for each activity using input and local variables. An activity

Table 1
Decomposition operators [28].

Operator	Description
optor_seq	Zero or more of the activities or actions in the decomposition must execute in any order, one at a time.
optor_par	Zero or more of the activities or actions in the decomposition must execute in any order and can execute in parallel.
or_seq	One or more of the activities or actions in the decomposition must execute in any order one at a time.
or_par	One or more of the activities or actions in the decomposition must execute in any order and can execute in parallel.
and_seq	All of the activities or actions in the decomposition must execute in any order, one at a time.
and_par	All of the activities or actions in the decomposition must execute in any order and can execute in parallel.
xor	Exactly one activity or action in the decomposition executes.
ord	All activities or actions must execute in the order (left to right) they appear in the decomposition.

can have three different conditions asserting what must be true for the activity to start executing (*Precondition*); repeat (*RepeatCondition*), or complete (*CompletionCondition*).

A decomposition has an operator that constrains how many of the acts in that decomposition can execute as well as the (ordinal) temporal relationships between them. EOFM supports ten such operators, those topical to the paper are shown in Table 1.

Atomic human actions or internal actions (representing cognitive or perceptual behaviors) occur at the bottom of every task. Human actions can have one of three behaviors: *AutoReset* actions occur as a single event; *Toggle* actions switch between occurring and not occurring; and *SetValue* actions commit some value to the other parts of the system. Internal actions allow cognitive behaviors (such as remembering something) to be modeled by assigning values to local variable.

To enable unambiguous interpretation of EOFM behavior, EOFMs have formal semantics [23,31]. In this, every act is interpreted as a state machine (Fig. 1) with three states: *Ready*, *Executing*, and *Done*. All acts start in *Ready*. They transition between states based on whether the Boolean conditions on the labeled transitions (Fig. 1) are true. When an action is *Executing*, this corresponds to an output variable associated with the action being set to the appropriate value. For *AutoReset* actions this is a Boolean variable that is set to true. For *Toggle*, the variable's Boolean variable is set to the negation of its current value. For *SetValue* behavior, the action's output variable is set to the value assigned in the model. The same occurs when the action is a local variable assignment.

Any existent activity strategic knowledge conditions (*Preconditions*, *RepeatConditions*, and *CompletionConditions*) partially describe when transitions can occur. However, there are three additional implicit conditions (based on the activity's or action's task position) that also impact transitions.

The *StartCondition* indicates if an act can start executing based on the states of its parent and siblings (acts in the same decomposition) as well as the parent's decomposition operator. Mathematically, a *StartCondition* has the generic form:

$$\text{StartCondition} : \text{parent.state} = \text{Executing} \wedge \bigwedge_{\forall \text{siblings } s} (s.\text{state} \neq \text{Executing}). \quad (1)$$

If the parent's decomposition allows for parallelism (the operator ends with *_par*), the second conjunct is eliminated. If the parent has an *ord* decomposition, to enforce order, the second conjunct requires that the previous sibling in the decomposition be done: (*prev_sibling.state* = *Done*). If the parent has an *xor* decomposition, the second conjunct is modified so that no other sibling can execute after one has finished: ($\bigwedge_{\forall \text{siblings } s} (s.\text{state} = \text{Ready})$). An activity without a parent (a top-level activity) will eliminate the first conjunct. Top-level activities treat each other as siblings in an *and_seq* formulation for the second conjunct.

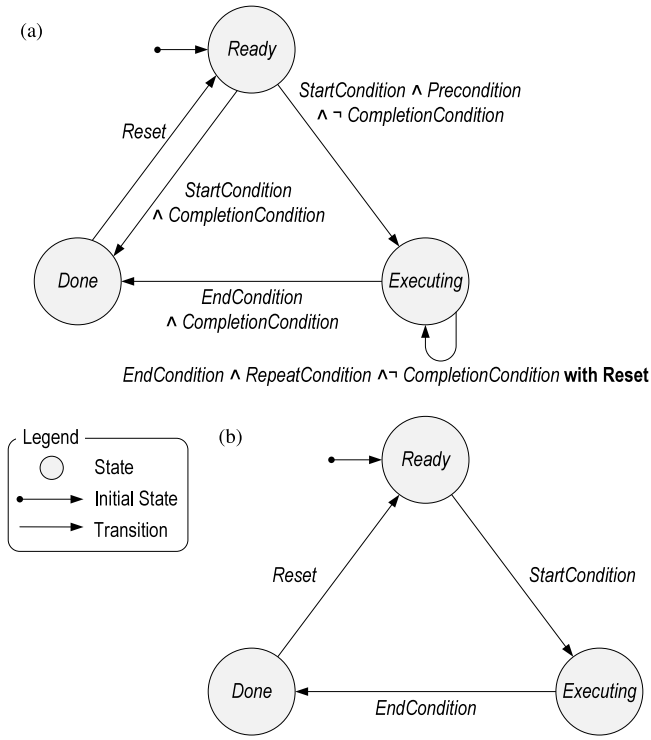


Fig. 1. EOFM transition semantics. (a) Transition semantics for an EOFM activity. (b) Transition semantics for an EOFM action.

The *EndCondition* indicates if an act can stop executing based on the execution state of its children¹:

$$EndCondition : \bigodot_{\forall subacts\ c} (c.state = Done) \wedge \bigwedge_{\forall subacts\ c} (c.state \neq Executing). \quad (2)$$

The first conjunct requires that the execution states of the activity's children satisfy the activity's decomposition operator. Here, $\bigodot$ is generic and substituted with $\bigwedge$ if the activity has an *and_seq*, *and_par*, or *ord* decomposition. It is substituted with $\bigvee$ for activities with *or_seq*, *or_par*, or *xor* decompositions. Because *optor_seq* and *optor_par* do not impose restrictions on the number of children that can execute, the first conjunct is eliminated for either decomposition. The second conjunct always specifies that no children are *Executing*.

The *Reset* describes when an activity can return to *Ready*. It occurs in two situations. First, when an activity repeats (a normative *Executing-to-Executing* transition), every descendant act *Resets*. Second, any root activity automatically *Resets* from *Done*. When this occurs, all descendants *Reset*.

The formal semantics of EOFM were used to create automated translation software that converted EOFM XML into the input language of the Symbolic Analysis Laboratory (SAL), a suite of non-probabilistic model checkers [12,23,31]. This enables EOFM behavior to be used as part of larger formal system analyses. As a result, EOFM analyses have been used to assess the impact of normative human behavior on a number of automotive, aerospace, and medical systems [22,32–34]. EOFM has also been used to model and predict the impact of erroneous human behavior.

2.2.2. EOFM and erroneous human behavior

EOFM has been used to automatically generate human errors using multiple theories [12,24–26,28,35]. This culminated in EOFM serving as the basis for the task-based taxonomy of erroneous human

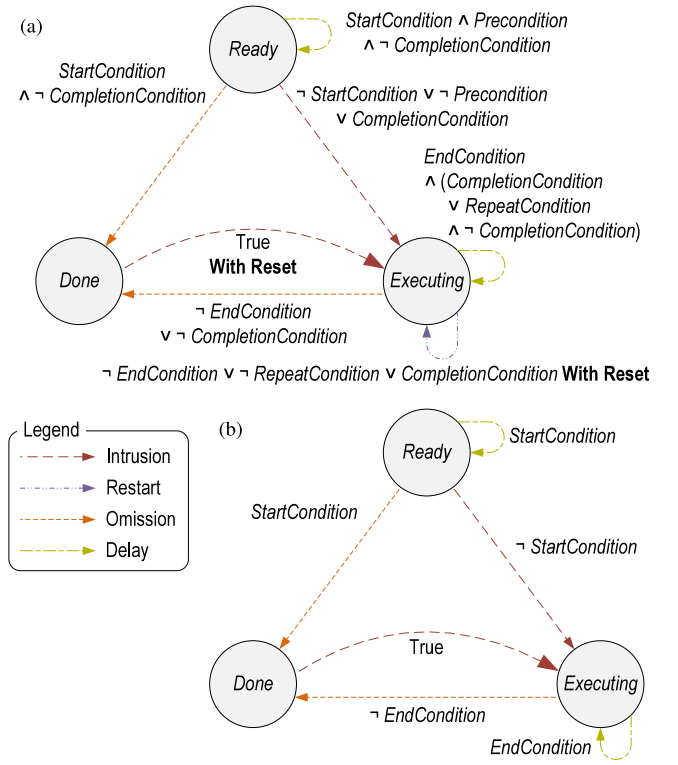


Fig. 2. Task-based taxonomy of human error transition-based error modes.

behavior [36] and the associated error generation approach [29]. The task-based taxonomy unifies the phenomenological (what) and genotypical (why) perspectives by classifying human error based on where they deviate from tasks. By knowing where a deviation occurs in a task model (how the task's formal semantics were violated), one understands how the error will observably manifest (the phenotype; [37]) and why the error occurred (the strategic knowledge or inherent condition the human improperly attended to; the genotype of the slip; [38]). In fact, the task-based taxonomy has been shown to be complete with respect to both the leading phenomenological [37] and genotypical [38] taxonomies [36].

The task-based taxonomy has a hierarchical classification that goes beyond what is relevant to this discussion [36]. What is topical is that it distinguishes between semantic violations that are transition-based (transitions that violate what is in Fig. 1) and assignment-based (incorrect behavior assignment to an action variable during action execution for *SetValue* or local variable actions). Transition-based error modes describe erroneous transitions between act execution states (see Fig. 2). An intrusion happens when an act executes (transitions to *Executing*) when it should not. An omission manifests when an activity finishes (transitions to *Done*) when it should not. A restart occurs when an activity (but not an action) incorrectly restarts: the activity resets and starts executing anew. Finally, a delay happens when an act does not transition out of a state when it should. In the execution-based erroneous behaviors, there can be both value and target substitutions (for *SetValue* actions) and misremembrances (for local variable assignments) (Table 2). In a value substitution or misremembrance, an incorrect value is assigned to the correct target variable (the action's output or local variable). In a target error, the correct value is assigned to the wrong target (the wrong output or local variable).

This taxonomy was the basis for a method [29] to automatically generate erroneous behavior in model checking analyses to potentially discover (and fix) human errors that could cause system failures. Specifically, this method would modify the EOFM-to-SAL translator so that all erroneous transitions (Fig. 2) and assignments (Table 2)

¹ Because actions have no children, action *EndConditions* default to true.

Table 2
Action-level erroneous behavior types.

Assignment	Erroneous behavior type
<i>CorrectAction</i> := <i>IncorrectValue</i>	Action Value Substitution
<i>IncorrectAction</i> := <i>CorrectValue</i>	Action Target Substitution
<i>CorrectVariable</i> := <i>IncorrectValue</i>	Value Misremembrance
<i>IncorrectVariable</i> := <i>CorrectValue</i>	Target Misremembrance

could exist alongside their normative counterparts (Fig. 1). A maximum (*Max*) on the total number of errors included was used to keep models from becoming completely unbounded. The net effect of this was that analyses accounted for all of the different ways that *Max* errors could occur when evaluating system safety.

This approach is powerful and particularly good at discovering HAI design flaws that could interact with unanticipated erroneous human behaviors. However, it does have limits. In particular, it does not account for the relative likelihood of different erroneous behaviors. This means that analyses may discover errors that are extremely unlikely to occur. Thus, with no means of ranking the relative probability of different errors, it can be difficult for analysts to prioritize how to address discovered problems. In fact, none of the formal verification work that has used task analytic or cognitive human behavior has accounted for the probabilities of different errors. To address this shortcoming, we plan to use concepts from HRA.

2.3. HRA and CREAM

HRAs are used to predict human error rates. There are many different HRAs [39]. These generally fall into two categories: those based on probabilistic risk assessment (the first generation) and those based on cognition (the second generation) [40]. Third generation techniques (which are discussed subsequently) exist, but use the theoretical foundation of first- and second-generation methods.

First-generation methods like the Technique for Human Error Rate Prediction (THERP) [41] and the Human Error Assessment and Reduction Technique (HEART) [42] (among others) are useful. However, they generally treat human errors the same as equipment failures. This limits their applicability because they do not account for the effect of context and organizational factors [40,43]. Second-generation HRAs build off and improve on these by considering the interactions inherent in complex systems (between humans, processes, organization, and the environment) based on how they affect human cognition [43–45]. CREAM is largely the preferred HRA by human factors engineers because it is rooted in a well-established cognitive model and can thus explain why errors occur [40,46–49]. Beyond this, CREAM has been well validated over the years through its successful use in nuclear power plants [50], manufacturing [51], radiation therapy [52], hospitals [53], and many other critical domains [43,54–60].

In CREAM [43], analysts identify the major tasks for working with a given system. The analysts then performs subjective assessments with subject matter experts. First, common performance conditions (or CPCs) are assessed for each task to understand how the work environment impacts the performance of that task. This means determining if each of the nine CPCs (quality of the organization, work conditions, human-machine support, procedures, simultaneous goals, time availability, time of day, work experience, and team collaboration) improves, reduces, or does not impact human performance. By synthesizing these ratings (counting the number that improve and reduce performance, and adjusting for dependencies), an analyst determines if the person is operating at one of the four Contextual Control Model (COCOM) modes (strategic, tactical, opportunistic, and scrambled; listed in decreasing levels of human reliability). In the basic version of CREAM, these control modes map to a range of probabilities of error occurring.

Table 3
CREAM's Cognitive function failures [43].

Cognitive function	Cognitive function failure
Observation	Wrong object observed Wrong identification Observation not made
Interpretation	Faulty diagnosis Decision error Delayed interpretation
Planning	Priority error Inadequate plan
Execution	Action of wrong type Action of wrong time Action on wrong object Action out of sequence Missed action

Table 4
Parameter values used for calculating probabilities of human error in the revised version of CREAM [61].

Parameter	Cognitive function			
	Observation	Interpretation	Planning	Execution
<i>C</i>	−2.0775	−1.3495	−2.0000	−2.4120
<i>a</i>	0.0055	0.0041	0.0052	0.0065
<i>b</i>	−0.2458	−0.2046	−0.2828	−0.2860
<i>c</i>	0.2840	0.2244	0.4019	0.4079

Point estimates can be achieved with two versions of extended CREAM. In these, an analyst identifies each task's cognitive function: observation — observing information in the environment; interpretation — understanding observed information; planning — setting a course of action; and execution — performing the planned actions. Next, the analyst identifies the function's most likely cognitive function failure (Table 3). Each failure has an associated nominal probability, originally identified by analyzing a large database of human performance information [43]. In the first approach to extended CREAM, the nominal probability is scaled to account for the impact of CPCs by multiplying it by a scaling factor associated with the assessed control mode. In the second variant, the nominal probability is scaled based on the impact specific CPC levels have on the different cognitive functions.

CREAM has been used successfully in a number of safety critical environments. However, it does have some problems. Bedford et al. [61] noted large error ranges and inconsistencies between predictions made with the three versions of CREAM. To address this, they created an improved version that reconciled the basic and extended methods (itself a refinement of previous attempts to improve CREAM predictions; [44,62]). First, the method accounts for the impact of each CPC on overall performance in a single integer, thus avoiding the collapse of ratings into course control mode designations. This is computed as the total number of CPCs that improve human performance and the number that reduce it:

$$CPC_{Sum} = \sum_{cpc \in CPCs} \begin{cases} 1, & \text{if } cpc \text{ improves} \\ 0, & \text{otherwise} \end{cases} - \sum_{cpc \in CPCs} \begin{cases} 1, & \text{if } cpc \text{ reduces} \\ 0, & \text{otherwise} \end{cases} \quad (3)$$

Second, using the same data that was the basis for CREAM's predictions [43], Bedford et al. [61] fitted a regression model for predicting the probability of human error

$$P_{HumanError} = 10^{(C+a \cdot CPC_{Sum}^2+b \cdot CPC_{Sum}+c)} \quad (4)$$

where *C* is the log₁₀ of the nominal probability of error for the cognitive function and *a*, *b*, and *c* are the regression coefficients for that function.

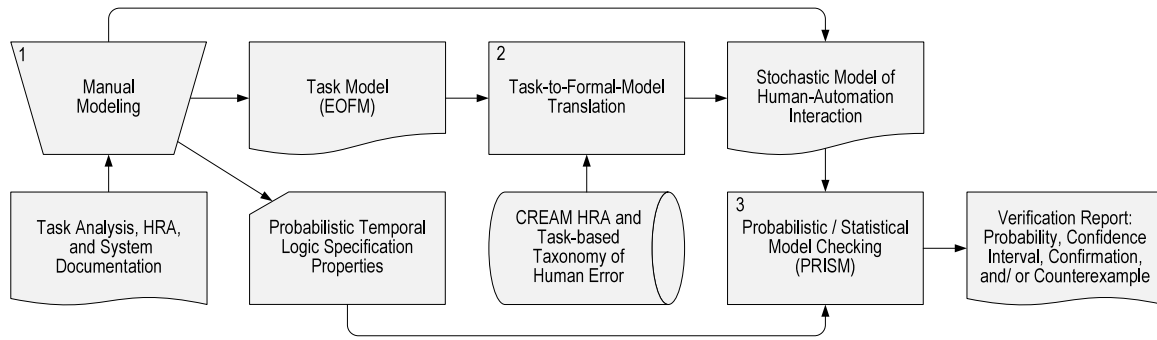


Fig. 3. Flow diagram of our formal method for including probabilistic erroneous behavior in formal verification. Numbers in processes show order of operation.

See Table 4 for the computed values. Thus, by computing the CPC_{Sum} with (3) and using the result in (4) with the parameters from Table 4, accurate predictions of human error probabilities for a given task can be computed.

While CREAM (and its variants) provides a foundation for predicting error probabilities, it is a manual process that provides no mechanism for analyzing whether task errors will actually be allowed to occur and whether they will impact system safety and performance. Combining CREAM with formal methods or simulation allows these issues to be addressed in what are called third generation approaches.

2.4. Third generation HRA

Third generation HRAs implement first and second generation methods in simulation or formal analyses to account for dynamic system behavior in predictions. For example, SHERPA (a simulator for human error probability analysis; [63]), like most third generation HRAs [40], is based in first generation methods and uses simulation. By being first-generation-based, methods like SHERPA are limited by not including a cognitive model. Their simulation basis also means that it can miss critical interactions that would be considered with formal methods. To address these problems, the Systems Analysis for Formal Pharmaceutical Human Reliability (SAFPHR) [57,59,60] showed that basic [57,59] and extended [60] CREAM could be implemented in the PRISM probabilistic model checker [64]. When this includes a dynamic model of the system, SAFPHR was able to accurately predict the probability of undesirable outcomes in community pharmacies. These developments are powerful and made several specific recommendations that could improve community pharmacy [60]. However, they are limited in two important ways. First, SAFPHR uses a very simple version of task analysis that is based on flow diagrams and thus does not allow for the specificity of behavior supported by EOFM. Second, SAFPHR requires analysts to manually identify the cognitive function failure for each task and thus does not dynamically determine how errors could occur.

3. Objectives

In this work, we hypothesized that the expressive power of the task-based taxonomy of human error [36] would give us enough information to associate every one of the possible errors with a cognitive function from CREAM (Table 3). This would enable us to automatically generate erroneous human behavior as part of a formal model (as was done in [29]) but with an associated probability of error for each possible task deviation. In accomplishing this, we would be able to use probabilistic and/or statistical model checking to compute the probability of system safety and performance specifications being violated and for understanding the likelihood of different errors contributing to failures. In what follows, we describe how a method based on these capabilities was realized. This is followed by a demonstration/validation of our approach by applying it to a simple but realistic example, an ATM. In doing this, we were able to obtain preliminary results of the methods

scalability and compare probabilistic and statistical model checking approaches. After this, we discuss the import of our findings and how they could influence future research.

4. Method

Our method for stochastically modeling erroneous human behavior in formal verifications is shown in Fig. 3. In this, an analyst has access to system information, a completed task analysis, and HRA (CPC_{Sum} s for each task). The analysts uses this to create a task model using PEOFM (which incorporates HRA information). He or she also creates part of a formal system model that describes the behavior of the system the human interacts with and specification properties that he or she wishes to check. Next, the analyst uses a systematic translation process that employs theory from CREAM and the task-based taxonomy of human error to incorporate the task model into the formal system model. This accounts for erroneous human behaviors and their probabilities. Then, the analysts run either probabilistic or statistical model checking (using PRISM) to verify specification properties against the formal model. This produces a verification report that, depending on the property checked and checking method, will produce a probability, CI around a probability, confirmation, or counterexample.

The three major contributions of this method relate to the extensions to EOFM required for this method, the novel translation process, and the specification properties that can be checked. We describe each of these below.

4.1. EOFM extensions

To support the new method, changes were made to the EOFM language [23], now Probabilistic EOFM (PEOFM). First, PEOFM only supports eight of the EOFM decompositions (those from Table 1). The sync (where decomposed actions are all performed synchronously) and com (for communicating information between humans) operators were removed because it was not clear how they would work with the probabilistic features. This issue is further explored in Section 6. Second, an optional $cpcsum$ attribute was added to the base element (eofm) of each task (Fig. 4(a)). This was designed to set the CPC_{sum} from (3). Third, because nondeterminism in PRISM model initial values limits analyses, an initial value was added to $SetValue$ actions (Fig. 4(b)).

4.2. Translation

The translator, like EOFM's non-probabilistic translators [23,29], interprets EOFM tasks as state machines using EOFM's normative (Fig. 1) and erroneous semantics Fig. 2 and Table 2. Older versions translated into the language of SAL [23,29] and human errors were non-probabilistic and limited by an analyst-specified maximum. The new

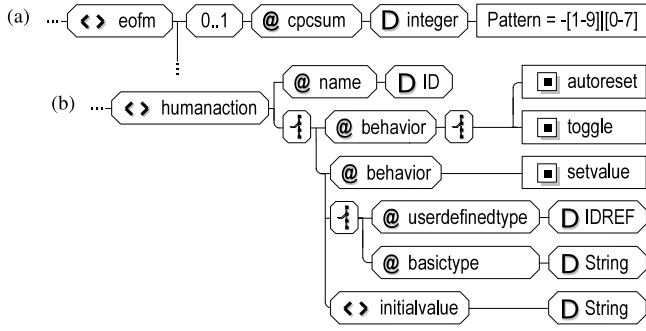


Fig. 4. A visualization [65] of additions made to EOFM to create PEOFM. (a) A cpcsum attribute for each task. (b) An initial value for SetValue actions. A ... indicates when the remainder of the language is the same as reported in [23].

Table 5
Mapping of CREAM cognitive functions to errors from EOFM's task-based taxonomy.

Cognitive function	Task-based human errors
Observation	Any action-level error related to local variable assignment, including target and value misremembrances
Interpretation	Any activity-level error (intrusion, omission, restart, or delay) caused by a violation of strategic knowledge conditions
Planning	Any activity-level error (intrusion, omission, restart, or delay) caused by violations of inherent or strategic knowledge conditions
Execution	Any action-level error that is not a local variable assignment

translator converts task models into the language of PRISM² and accounts for human error probabilities.

The key to including probabilities in translation is a mapping (Table 5) between CREAM's cognitive functions (Table 5) and the deviations from normative EOFM semantics from the task-based taxonomy of human error Fig. 2 and Table 2. Because local variable assignments are the mechanism EOFM uses for representing humans noticing and remembering things from the environment, errors associated with this represent observation errors. CREAM suggests that interpretation relates to how humans contextualize observed information. This maps to activity-level errors caused by the incorrect interpretation of strategic knowledge (activity *Preconditions*, *RepeatConditions*, and *CompletionConditions*; Fig. 2(a)). Planning errors in CREAM occur because people improperly form or understand a plan or task. Thus, map to any violation of the logic that people need to understand to properly execute the task: any activity-level transition error (Fig. 2(a)) that can be caused by violations of inherent or strategic knowledge conditions. According to CREAM, execution errors occur when the person is attempting to physically execute task actions. In our mapping, execution errors all occur at the action level, either through the transition of action execution states (Fig. 2(b)) or in the assignment of an output value (the first two rows of Table 2).

Translation was implemented as a Java desktop application that converts a PEOFM into PRISM's language as a Markov decision process using the architecture shown in Fig. 5. First, the formal EOFM representation is contained in a single module. This captures the behavior (normative and erroneous) associated with each task. This task module interacts with a module (or modules) that represent other parts of the system (Sys in Fig. 5). This can include automation behavior, the environment, or human mission parameters as in traditional EOFM verifications [22]. The interface between the human and their parts of the system is represented by shared variables between modules. Outputs from the human task are actions and outputs from the other elements represent display information, environmental conditions, or

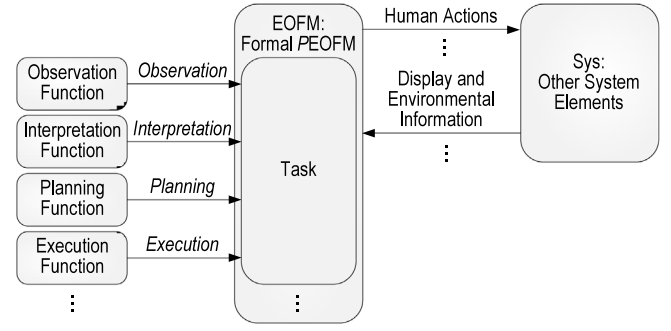


Fig. 5. Architecture used for representing a PEOFM in PRISM. Modules are rounded rectangles. Arrows are variables shared between modules with input (destination) and output (source) relationships.

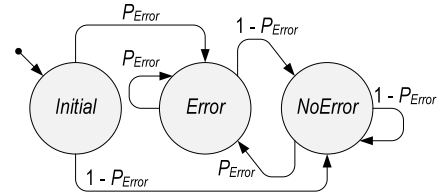


Fig. 6. Transition logic for the value/state of the output variable of cognitive function modules (see Fig. 5) at each model step. These transitions are probabilistic, where P_{Error} is calculated using (5).

mission perogatives. To account for the probabilities for performing tasks normatively or erroneously, the EOFM module also takes inputs from modules representing each CREAM cognitive function (Table 3). There are four such modules for each EOFM task. The probabilistic behavior of the model is implemented in these cognitive function modules. Specifically, each such module produces an output that is sent to its associated task that indicates if that function occurs with an error. Thus, at each model step, the EOFM's task model will select from the set of available behaviors based on the execution state of the task and its cognitive functions.

A task's cognitive functions compute the probability of an error (using (4) with the appropriate parameters from Table 4) as

$$P_{Error} = \begin{cases} \min \left(10^{(C+a \cdot CPC_{Sum}^2 + b \cdot CPC_{Sum} + c)}, 1 \right), & \text{if } CPC_{Sum} \geq -9 \\ 10^C, & \text{otherwise.} \end{cases} \quad (5)$$

This differs slightly from (4). Because (4) does not account for situations where the equation produces probabilities greater than one, the first case of (5) selects the minimum of (4) and 1. Additionally, PEOFM markup allows CPC_{Sum} to go undefined (see Section 4.1). In this situation, the translator assumes a default of CPC_{Sum} that is less than -9. When this occurs, the second case of (5) assumes the cognitive function's nominal error probability [43,61].

Cognitive function outputs all behave as shown in Fig. 6. All start in an *Initial* state. For any given state, the output will indicate *Error* with probability P_{Error} and *NoError* with probability $1 - P_{Error}$ in the next state. The *Initial* value allows cognitive functions to transition to *Error* or *NoError* with the appropriate probabilities before EOFM module transitions (which all require *Error* or *NoError* values) can occur.

While considering the output of the cognitive function modules, the translator interprets the semantics from Figs. 1 and 2. It does this by creating a variable in the EOFM module representing the execution state of each act and explicitly representing the transitions in Fig. 7(a) and (c) for activities and Fig. 7(b), (d), and (e) for actions. These implement the original semantics except they only allow transitions to occur if the behavior is enabled by the task's cognitive function outputs.

² <https://www.prismmodelchecker.org/manual/>

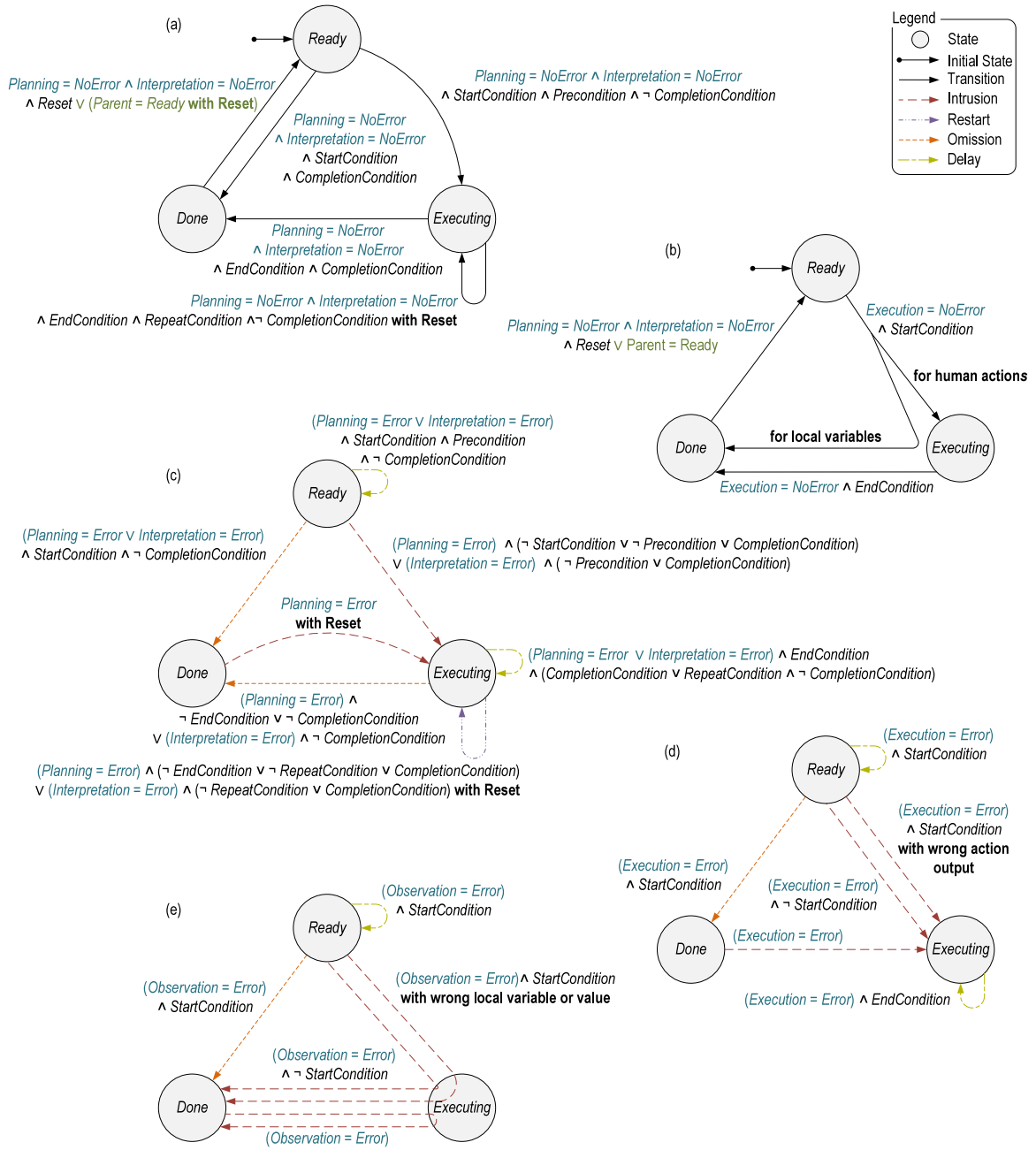


Fig. 7. Transition diagram showing how the PEOFM to PRISM translator interprets the transition semantics from Figs. 1 and 2. Color is used in transition logic to highlight new conditions. Note that here *Observation*, *Interpretation*, *Planning*, and *Execution* represent the output variables from the associated task's cognitive function modules (Figs. 5 and 6). (a) How the translator interprets normative transitions for activities from Fig. 1(a). (b) How the translator interprets normative transitions for actions from Fig. 1(b). (c) How it interprets erroneous transitions for activities from Fig. 2(a). (d) How it interprets erroneous transitions for actions that execute human actions from Fig. 2(b). (e) How it interprets erroneous transitions for actions that execute a local variable assignment.

So, for example, the transitions in Fig. 7(a) and (b) represent normative behavior (from Fig. 1) and can thus only occur when there are no errors in the associated cognitive functions. In the transitions, there is an additional new condition on *Reset* that allows it to occur when the act's parent is ready (*Parent = Ready*). This gives an act the option to reset after its intrusion has completed. Consistent with previous non-probabilistic translators, actions with local variable assignments automatically transition from *Ready* to *Done*, with the variable assignment occurring in the process. This is done because modules external to EOFM cannot observe this action and this form of the transition improves model scalability.

Fig. 7(c)–(e) represent the erroneous behaviors from Fig. 2. Additional conditions are added to account for the mappings for observation, interpretation, planning, and execution errors from Table 5. Note that (d) and (e) describe the difference ways that errors manifest for human actions (which rely on the execution cognitive function) or local variable assignments (which depend on observation) respectively. In both of these, the *Executing*-to-*Done* transition is eliminated. This is because an action's *EndCondition* is true by default when the action is performed and action performance time is not currently modeled (actions are treated as being instantaneous). This means that an *Executing*-to-*Done* omission is functionally equivalent to a *Ready*-to-*Done* omission. Also note that in (e), *Ready*-to-*Executing* and *Done*-to-*Executing* intrusions skip the *Executing* state for the same reason that

normative local variable assignments do. Finally, CREAM allows for incorrect actions to be performed (for the execution function) or the wrong object observed (for the observation one) (Table 3) as a single error for the respective function. The original erroneous transitions (Fig. 2) can support this behavior, but require both an intrusion and omissions. Thus, to represent wrong action and wrong object observed error types as a single error in our method, both (d) and (e) contain an extra transition (*Ready-to-Executing* in (d) and *Ready-to-Done* with implied execution in (e)) to allow an incorrect action or incorrect local variable assignment (respectively) to be performed when the given action should execute normatively.

As in previous translators [23,29], the new one creates a variable representing the output or product of each human action. For *AutoReset* and *Toggle* actions, these are Boolean. For *SetValue* actions and local variable assignments, the variables have the type specified in the PEOFM markup. When the normative action transitions from *Ready* to *Executing* (or from *Ready* to *Executing* with presumed execution for local variable assignments; Fig. 7(b)), the corresponding output or local variable is set to its appropriate value: true for *AutoReset*, the negation of its current value for *Toggle*, and the markup-specified value for *SetValue* and local variable assignments.

Erroneous action intrusions also exhibit this behavior, but vary in terms of what variable is assigned or what value is assigned to it. For all actions, the intrusions that transition from *Ready* with a negated start condition have the same target variable and value that would be assigned normatively. This is also true of intrusions that originate from the *Done* state. Both transitions represent a situation where an action is inserted (extraneously) into the executing task. Action intrusions originating from *Ready* with a normative (non-negated) *StartCondition* represent intrusions where the normative action is replaced with something else. Since this “something else” could be many different actions, the translator creates multiple transitions of this type. For a given human action, this means that the translator creates transitions where the human sets the output of every other possible human action instead of the one for the current action (one transition per other action). Similarly, for local variables, this means that the translator creates transitions where all of the other local variable assignments from the markup occur instead of the correct one. Additionally, these types of transitions are used for representing Action Value Substitutions and action target substitutions (for *SetValue* human actions) and value misremembrances and target misremembrances (for local variable assignments) (Table 2). For both action value substitutions and misremembrances, the translator creates multiple transitions where the value assigned to the correct variable is wrong in each possible way that it could be wrong, one transition for every possible wrong value. For action target substitutions and target misremembrances, the translator creates transitions where the incorrect variable is assigned the correct value for every possible human action output or local variable (respectively) that has the same type as the correct variable.

Analysts manually complete the sys module (Fig. 5). However, the translator creates a template for this that defines the output variables and provides a pattern for module transitions.

4.3. Specification properties

Finally, analysts must create specifications to verify against models. Currently, these must also be manually created by the analyst. While many properties are possible, current efforts use probabilistic temporal logic specifications of the form:

$$P = ?[F(\text{FailureCondition})]. \quad (6)$$

When this specification is checked, it instructs the model checker to compute the probability ($P = ?$) that eventually (F) a failure condition (*FailureCondition*) occurs.

5. Application

As a proof of concept, we used the method to evaluate an ATM. This application was chosen because ATMs have well-documented design variations that can profoundly impact user errors. Specifically, some ATM designs can cause a postcompletion error (where a human omits actions related to subsidiary goals once the primary goal is achieved [66]): a user leaving the card in the machine after receiving cash. There are also known probabilities for postcompletion errors [67,68], enabling us to determine if our approach produces valid predictions. In what follows, we describe the behavior of two different ATM variants and how they were formally modeled. We also describe the task behavior the human operator uses when interacting with the ATM. We then describe how these models were integrated into complete formal models and used to evaluate the probability errors with the system.

5.1. The formal model of the ATM

Fig. 8 shows the formal models describing the behavior of the two variations of the ATM. In the first (Fig. 8(a)), the machine starts on a welcome screen. If the user enters a card, he or she is brought to the state for entering the pin. In this state, the user must enter his or her correct pin to progress, with incorrect pins keeping them in the current state. If the correct pin is entered, the user reaches the withdraw state. Here, the user must enter a monetary value that is greater than zero, then the machine outputs the cash, the user takes it, the machine then outputs the card. When the user takes the card, the machine returns to the welcome screen. When in the pin or withdraw states, the user can press a cancel button to immediately output the card, which can be retrieved by the user to return the machine to the welcome screen. The second machine (Fig. 8(b)) works the same as the first except that the order in which the cash and card are output is reversed. Note that the machine shown in (a) is the one that encourages postcompletion error because it allows the human to take the cash (the human’s primary goal) before taking the card (a subsidiary goal).

5.2. The task model for interacting with the ATM

Fig. 9 visualizes the PEOFM representing the human task behavior for withdrawing cash from both machines in Fig. 8.³ In this, a human who wants to withdraw cash (*aGetCash*) must execute four sub-activities in order (as dictated by the ord decomposition). The user first inserts his or her card (with the *EnterCard* *AutoReset* action via the *alInsertCard* activity) at the welcome screen and will continue trying to do so until this screen has cleared (as dictated by the activity’s strategic knowledge). The user will enter his or her correct pin (via the *EnterPin* *SetValue* action) while on the *EnterPin* interface state under the *aEnterPin* activity. Next, the user enters the desired cash value while in the *Withdraw* interface state. Finally, the user will retrieve machine outputs (*aRetrieveOutputs*) by performing two activities in any possible order (an *and_par* decomposition) based on when the system enables them: taking the cash and taking the card. The user knows that the activity (*aRetrieveOutputs*) is completed (via a completion condition) if the user thinks that the cash has been retrieved.

For the purpose of this analysis, we assume a $CPC_{Sum} = 4$ for the task. This was done because it represents the minimum value that

³ Task model XML as well as all model code used for this paper can be found at <http://fhsl.eng.buffalo.edu/resources/ProbabilisticEOM/>.

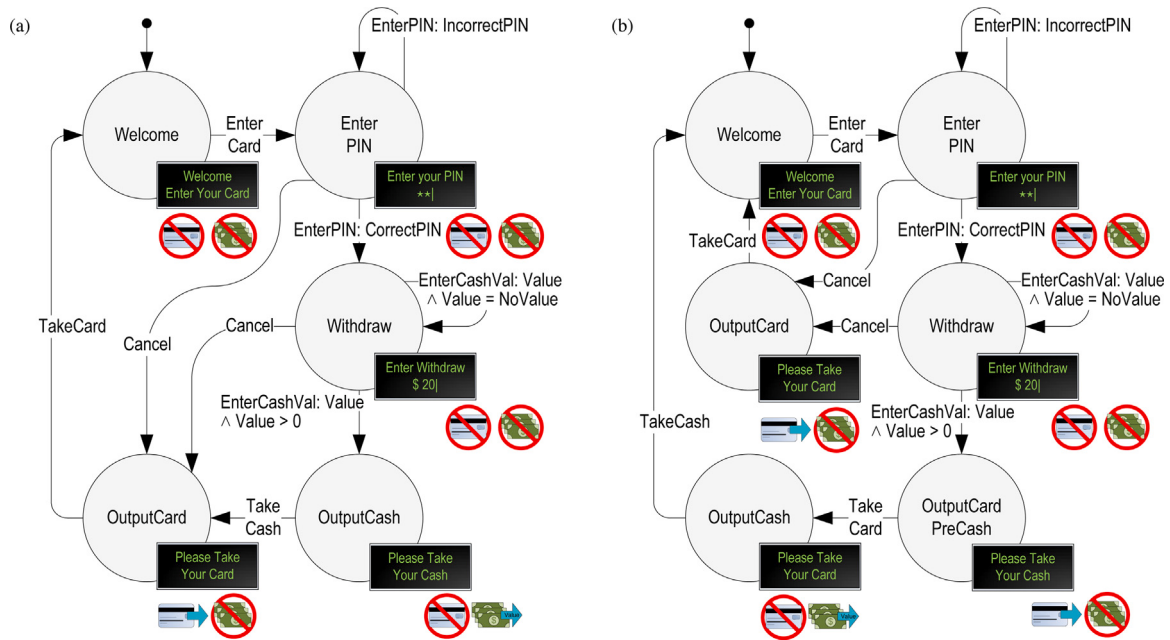


Fig. 8. State machine representation of the two variations of the ATM formal model. (a) An ATM that outputs cash before outputting the user's card. (b) An ATM that outputs the user's card before outputting cash. For both figures, each state encompasses the state of the interface (*iInterface*; which assumes the value shown in each circle), whether or not a card is being output (*iCardPresent*; a Boolean variable that is false except in the *OutputCard* state), and the state of cash output (*iCashout*; which defaults to *NoValue* but will assume the value entered by the user in the *OutputCash* state).

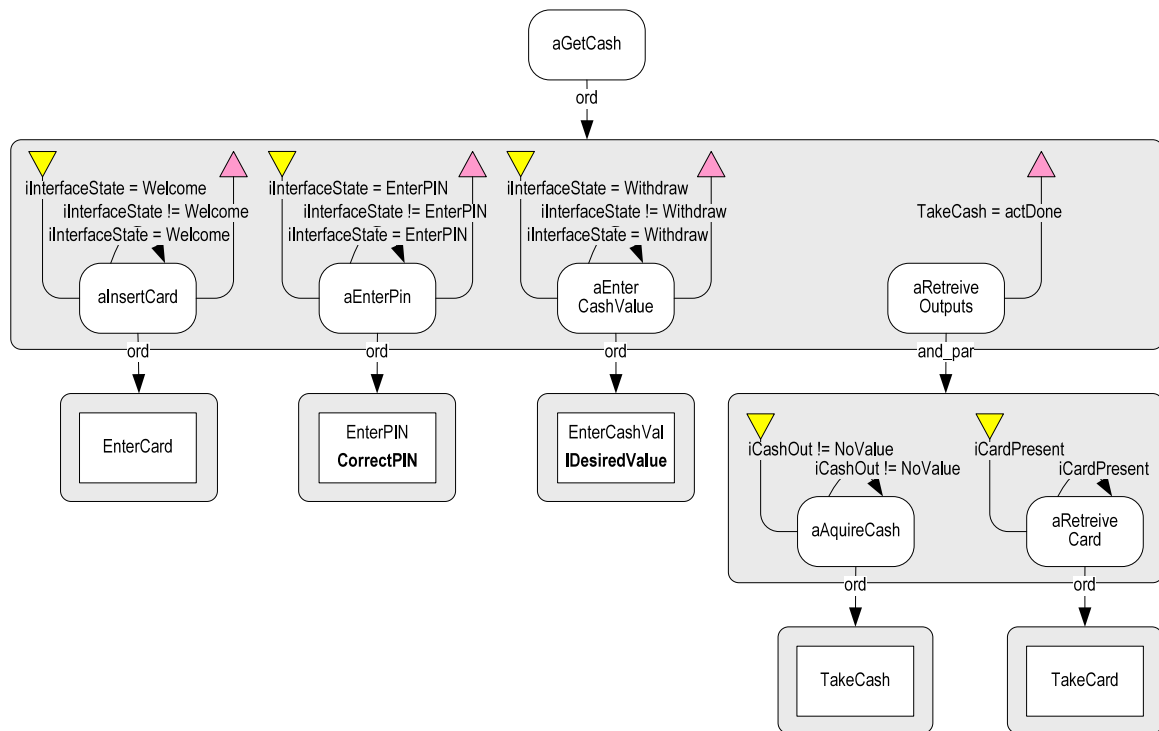


Fig. 9. Visualization of the PEOFM a human uses to interact with the ATMs from Fig. 8. This uses EOFM's visual notation [69]. Activities are rounded rectangles and actions are pointed rectangles. Activity decompositions are presented as rounded rectangles below an activity that are connected by an arrow labeled with a decomposition operator. *SetValue* actions are presented with both the actions name at the top and the value being committed below it **bolded**. Strategic knowledge conditions are labels on lines and arrows connected to their associated activities. Preconditions are down-pointing yellow arrows connected by on the left of the activity. Completion conditions are up-pointing magenta triangles connected on the right. Repeat conditions are recursive arrows on the top of the activity.

would normally be associated with the strategic control mode, implying that the human has a deep understanding of the task he or she is performing and can behave strategically. Given that most people are very familiar with ATMs and the way they work, this seemed like a reasonable assumption.

5.3. Translation and formal system model construction

The XML of the PEOFM from Fig. 9 (which was 76 lines long) was converted into the input language of the PRISM model checker, producing a representation that is 834 lines. This Sys module was

then completed in this translated version in two ways: one using the behavior from Fig. 8(a) and one using the behavior from Fig. 8(b). Thus we ultimately had two formal system models for comparison: one where we would expect a higher probability for having a credit card left in it (the one based on Fig. 8(a)) and one where we would expect this to be lower (the one based on Fig. 8(b)).

5.4. Specification properties

We evaluate our model with three specifications, all implemented using the pattern from Eq. (6). The first checked for the presence of the postcompletion error outcome: the person completing the task having received cash but leaving the card:

Card Left :

$$P = ? \left[F \left(\begin{array}{l} (aGetCash = actDone) \\ \wedge \\ iCardPresent \\ \wedge \\ \neg iCashOut = NoValue \end{array} \right) \right]. \quad (7)$$

This tells the model checker to calculate the probability ($P = ?$) that eventually (F) the task completes ($aGetCash = actDone$) with the cash having been retrieved (no longer being output by the machine; $iCashOut = NoValue$) and the card remaining as an output of the machine ($iCardPresent$).

The second checked for the person completing the task while leaving cash in the machine:

Cash Left :

$$P = ? \left[F \left(\begin{array}{l} (aGetCash = actDone) \\ \wedge \\ \neg iCardPresent \\ \wedge \\ \neg iCashOut > NoValue \end{array} \right) \right]. \quad (8)$$

Finally, overall reliability was assessed with a property requesting the probability of either (or both) errors occurring:

Reliability :

$$P = ? \left[F \left(\begin{array}{l} (aGetCash = actDone) \\ \wedge \\ iCardPresent \\ \wedge \\ \neg iCashOut > NoValue \end{array} \right) \right]. \quad (9)$$

If the method accurately predicts probabilities, we would expect the model containing the ATM from Fig. 8(a) to have a higher probability for leaving the card (7) than leaving cash (8) because of the machine's encouragement of the post-completion error. We would expect the probability of leaving the card to be lower for the model using Fig. 8(b) because this should not encourage the postcompletion error. We would also expect the probability of leaving the card to be similar to that of leaving the cash. Finally, we would expect the overall chance of error (9) to be higher for the model using Fig. 8(a) because of its support for the postcompletion error.

5.5. Verification and results

Verifications were performed on a workstation with a 12-core 3.60 GHz Intel Xeon E5-1650 with 128 Gigabytes of RAM. However, probabilistic model checking resulted in the machine running out of memory after more than 24 h of analysis. Luckily, we were able to use statistical model checking.

Statistical checking uses samples of model traces to compute CIs on computed probabilities. This means that scalability restrictions can be avoided at the expense of accuracy. To assess this tradeoff, we conducted a small experiment using the formal system model containing the machine behavior from Fig. 8(a) and checking the overall reliability with (9). This assumed a 99% CI computed from a sample size starting at 1000 and increasing in 1000 increments up to 25,000. Experiment results are shown in Fig. 10. This showed the verification time increased linearly with the number of samples Fig. 10(b). It also showed that

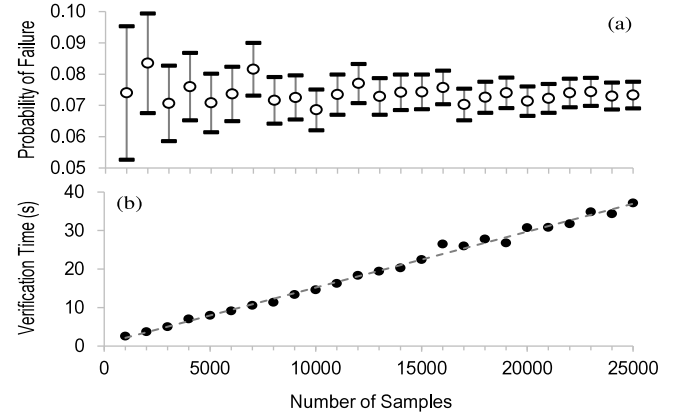


Fig. 10. (a) Verification results (predicted probability and its 99% CI) and (b) times from statistical model checking for verify (9) against models containing behavior from Fig. 8(a) for different numbers of samples.

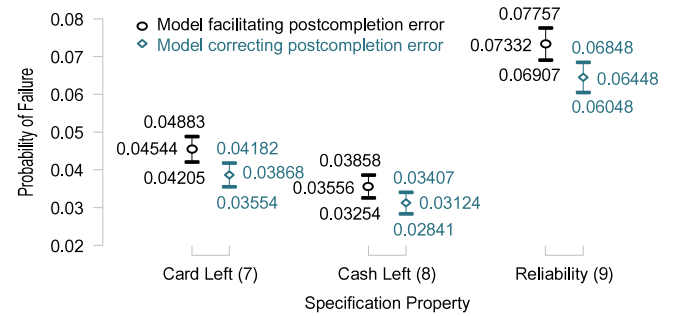


Fig. 11. Verification results (prediction and 99% CI) for checking the specifications in (7)–(9) against the two variations of the model: one with the behavior for enabling a postcompletion error (Fig. 8(a); black) and one with the behavior for avoiding it (Fig. 8(b); blue).

the width of the 99% CI narrowed as the sample size increased. At 25,000 samples, the verification took 37.146 s to estimate a probability of 0.07332 with a 99% CI of [0.069, 0.0776]. This interval is less than 0.01, or one percentage point. Thus, statistical model checking was able to produce results that we can be 99% confident are within 1 percentage point of the actual value in less than a minute of verification time. We used a sample size of 25,000 for subsequent verifications.

Fig. 11 shows the results of checking each of the properties against the two model variants. This produced results that are consistent with our expectations. The model with the behavior from Fig. 11(a) showed a probability of 0.04544 of leaving the card in the machine, a value significantly higher (shown by the CIs) than it was for the machine from Fig. 11(b). The second value was also comparable to the values seen for both models for leaving cash in the machine. Finally, the machine supporting the postcompletion error was significantly more likely to see any error (0.07332) than the other machine (0.06448).

6. Discussion

In this research, we have introduced a new formal method that combines task modeling, a taxonomy of erroneous human behavior, and HRA. This is able to both formally generate HAI errors and account for the probabilities of these errors. This enables engineers to evaluate reactive HAI designs by using probabilistic and statistical model checking to determine the relative probability of different outcomes both between and within designs. This is a major contribution for formal analyses of human error and reliability as previous approach could only determine if errors were possible. As such, this paper also makes significant contributions to reliability engineering and system

safety. Specifically, human behavior and human error play significant roles in failure of safety and reliability. The method introduced here give engineers an unprecedented ability to assess the reliability of safety-critical, human-interactive systems. Thus, this work will enable designs and interventions that will significantly reduce the probability of human error causing system failures and disasters, saving lives, and protecting critical infrastructure.

The ATM application is illustrative because it shows that probability predictions were able to differentiate between two designs in ways that match with established performance: the postcompletion error condition was significantly more likely with the interface that facilitated it than for the other interface and other errors. It is not entirely clear from the literature how likely an ATM postcompletion error is. However, Ratwani and Gregory Trafton [67], Ratwani and Trafton [68] in studies on postcompletion errors found that when they are facilitated by the HAI, they occur about 5% of the time. This 5% average is consistent with the 4.544% observed with our model. Thus, while limited, the ATM application does suggest that our approach is capable of producing valid predictions.

It is important to note that the ATM example produces more than just post-completion errors. This is evidenced by the fact that the overall error rates were higher than the post completion ones and that the model that did not allow for the post-completion error still exhibited errors. Of course, other systems will have different types of errors that will be critical to system safety and performance. Future research should seek to validate our method with other applications and human subject experiments. Additionally, there are extensions and considerations necessitated by our approach. We discuss these below.

6.1. Probability accuracy

While the presented results are consistent with the literature, the predictions seem high. This is not necessarily surprising because the method accounts for all of the possible ways an error can manifest, including all planning errors (incorrect task knowledge). For a task as familiar as interacting with an ATM, it is conceivable that planning errors might be less likely than for other tasks. It is worth noting that the method can be easily adapted to account for different conditions. First, researchers can assess the values of the CPCs used in a given environment and use those in analyses. Additionally, in situations where an analyst is sure a given cognitive function is not a factor, he or she can easily set the default probability of error in the associated function to 0 to mitigate the function's effect.

Additionally, the method enables analysts to evaluate the impact of different conditions on performance by manipulating the CPC values. For example, CREAM has a CPC representing the quality of conditions. Because any given CPC can assume a value of -1 , 0 , or 1 , it can have up to a 2 point impact on $CPCSum$. Any analyst wishing to see how environmental or design factors impact computed probabilities should be able to appropriately adjust CPC values and rerun analyses. This should be explored more thoroughly in subsequent research.

To some extent, the absolute accuracy of the probabilistic predictions is less important than relative accuracy. As long as the method can determine which condition is more likely, analysts will be able to take appropriate action. Future research will investigate the method's ability to predict probabilities both absolutely and relatively.

6.2. Deeper support for CPC variation

The current version of the method has a single $CPCSum$ and associated cognitive functions for each task. It is conceivable that CPC values could vary for different parts of a task. It should be possible for PEOFM to allow for greater specificity of $CPCSum$ values at the activity level and have the translator create separate cognitive functions to accommodate this. Future research should explore these developments.

6.3. Scalability

The ATM example showed that scalability is a major restriction of probabilistic model checking with PEOFM. The translation approach is derived from the original method that was used for including EOFM task behavior in nonprobabilistic model checking analyses [23]. In this, the execution state of every activity and action is represented with its own variable. In traditional model checking with EOFM, significant scalability improvement can be achieved by eliminating the explicit representation of EOFM activity execution state and expressing it as a formula over the execution state of actions [31]. This improvement works because it eliminates the need for intermediate task model transitions. However, this approach is currently incompatible with the human error generation method employed here because erroneous deviations occur at the intermediate transitions of the task's activities. It is, however, theoretically possible to apply the scalability improvement method from Bolton et al. [23] to PEOFM, but this would ultimately require a much more complex formulation. This is because there are an increased number of intermediary transitions necessitated by error generation and the need to account for the probabilities of these transitions. This should be explored more thoroughly in future research.

Fortunately, statistical model checking has fewer restrictions than probabilistic model checking. In fact, verification times appear to scale linearly with the number of computed samples. This provides good evidence that our method, when used with statistical checking, can scale to industrial applications. This should also be investigated further in future research.

6.4. Purely reactive system

It is important to note that our method is most appropriate for modeling reactive systems: where system behavior occurs in direct response to human actions. This is because the cognitive-function-based model assumes that stochastic behavior only originates with the human. There are many interactive systems that are purely reactive. However, this limits the applicability of our approach. To help expand the scope of applications that could be evaluated with the method, future research should investigate how to account for stochastic behavior originating from the environment and machine automated behavior.

6.5. Non-deterministic choice

A byproduct of using Markov decision processes and statistical model checking in our method is that nondeterministic transitions (where there are multiple allowable transition at a given step) occur with equal probability. For example, if there are two activities in an *and_seq* decomposition who both have their precondition satisfied, there will be a 0.5 probability of each executing. While this is normally the behavior an analyst will want, there are conceivable situations where humans would be more likely to choose one behavior over another. Future research should investigate how to account for this in the method.

6.6. The sync Decomposition and Human-human Interaction

PEOFM and the method currently support all but two of EOFMs features. First, PEOFM does not allow the sync decomposition, a rare condition where the human performs multiple actions synchronously. Second, the method does not support EOFMC capabilities (the variant supporting human-human coordination and communication; [12,30]). While CREAM's CPCs and cognitive functions do give some insight into how to account for human-human coordination errors, they do not appear to offer a theory for modeling the probability of communication errors. Future work should investigate how to incorporate EOFMC capacities into our method.

6.7. Dependence of error prediction on task modeling

Because our method uses the task-based taxonomy of human error, it is complete with respect to the phenotypes of erroneous action and slip genotypes [36]. This means that it should be capable of modeling nearly any type of human error and its associated probability. However, it is important to note that the type of errors that can be generated will depend of the way that the normative task model is formulated. For example, our method should be capable of generating errors where someone puts a card into the machine with incorrect orientation, or even inserts the wrong card. However, accomplishing this would require analysts to include actions in the task model for picking up different card options and inserting cards in different orientations. Doing this ultimately requires modeler insight that may not be obvious at the time of model creation. Future work should investigate how to create guidance for modelers that enables them to include model concepts that will enable complete error prediction. Additionally, progress has been made in using formal models of affordance to identify what human actions (intended or unintended) are facilitated by the interface and environment [70,71]. Future research should investigate whether affordance-based action prediction can be incorporated into our method.

CRedit authorship contribution statement

Matthew L. Bolton: Conceptualization, Methodology, Software, Formal analysis, Writing - original draft, Visualization. **Xi Zheng:** Conceptualization, Writing - review & editing. **Eunsuk Kang:** Conceptualization, Writing - review & editing.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgment

This material is based upon work supported by the National Science Foundation, USA under Grant Nos. 1918314 and 1918140.

References

- [1] Kohn LT, Corrigan J, Donaldson MS. To err is human: building a safer health system. Washington: National Academy Press; 2000.
- [2] Kenny DJ. 26th Joseph T. Nall report: General aviation accidents in 2014. Technical report, AOPA Foundation; 2015.
- [3] Kebabjian R. Accident statistics. 2018, <http://www.planecrashinfo.com/cause.htm>.
- [4] Manning SD, Rash CE, LeDuc PA, Noback RK, McKeon J. The role of human causal factors in US Army unmanned aerial vehicle accidents. Technical report 2004–11, USA Army Research Laboratory; 2004.
- [5] NHTSA. National motor vehicle crash causation survey: Report to congress. Springfield: National Highway Traffic Safety Administration; 2008.
- [6] Le Bot P. Human reliability data, human error and accident models—illustration through the Three Mile Island accident analysis. Reliab Eng Syst Saf 2004;83(2):153–67.
- [7] Bolton ML, Bass EJ, Siminiceanu RI. Using formal verification to evaluate human-automation interaction in safety critical systems, a review. IEEE Trans Syst Man Cybern Syst 2013;43(3):488–503.
- [8] Bolton ML. Novel developments in formal methods for human factors engineering. In: Proceedings of the human factors and ergonomics society annual meeting. Los Angeles, CA: SAGE Publications Sage CA; 2017, p. 715–7.
- [9] Weyers B, Bowen J, Dix A, Palanque P, editors. The handbook of formal methods in human-computer interaction. Berlin: Springer; 2017.
- [10] Paternò F, Santoro C. Integrating model checking and HCI tools to help designers verify user interface properties. In: Proceedings of the 7th international workshop on the design, specification, and verification of interactive systems. Berlin: Springer; 2001, p. 135–50.
- [11] Aït-Ameur Y, Baron M. Formal and experimental validation approaches in HCI systems design based on a shared event B model. Int J Softw Tools Technol Transf 2006;8(6):547–63.
- [12] Bolton ML, Bass EJ. Enhanced operator function model (EOFM): A task analytic modeling formalism for including human behavior in the verification of complex systems. In: Weyers B, Bowen J, Dix A, Palanque P, editors. The handbook of formal methods in human-computer interaction. Cham: Springer; 2017, p. 343–77.
- [13] Bastide R, Basnyat S. Error patterns: Systematic investigation of deviations in task models. In: Task models and diagrams for users interface design. Berlin: Springer; 2007, p. 109–21.
- [14] Fields RE. Analysis of erroneous actions in the design of critical systems [Ph.D. thesis], York: University of York; 2001.
- [15] Barbosa A, Paiva AC, Campos JC. Test case generation from mutated task models. In: Proceedings of the 3rd ACM SIGCHI symposium on engineering interactive computing systems. ACM; 2011, p. 175–84.
- [16] Clarke EM, Grumberg O, Peled DA. Model checking. Cambridge: MIT Press; 1999.
- [17] Kwiatkowska M, Norman G, Parker D. Stochastic model checking. In: Bernardo M, Hillston J, editors. Formal methods for the design of computer, communication and software systems: performance evaluation. LNCS (tutorial volume), vol. 4486, Berlin: Springer; 2007, p. 220–70.
- [18] Schraagen JM, Chipman SF, Shalin VL. Cognitive task analysis. Philadelphia: Lawrence Erlbaum Associates, Inc.; 2000.
- [19] Basnyat S, Palanque P, Schupp B, Wright P. Formal socio-technical barrier modelling for safety-critical interactive systems design. Saf Sci 2007;45(5):545–65.
- [20] Gunter EL, Yasmeen A, Gunter CA, Nguyen A. Specifying and analyzing workflows for automated identification and data capture. In: Proceedings of the 42nd hawaii international conference on system sciences. Los Alamitos: IEEE Computer Society; 2009, p. 1–11.
- [21] Palanque PA, Bastide R, Senges V. Validating interactive system design through the verification of formal task and system models. In: Proceedings of the IFIP TC2/WG2.7 working conference on engineering for human-computer interaction. London: Chapman and Hall; 1996, p. 189–212.
- [22] Bolton ML, Bass EJ. Formally verifying human-automation interaction as part of a system model: Limitations and tradeoffs. Innov Syst Softw Eng NASA J 2010;6(3):219–31.
- [23] Bolton ML, Siminiceanu RI, Bass EJ. A systematic approach to model checking human-automation interaction using task-analytic models. IEEE Trans Syst Man Cybern A 2011;41(5):961–76.
- [24] Bolton ML, Bass EJ, Siminiceanu RI. Generating phenotypical erroneous human behavior to evaluate human-automation interaction using model checking. Int J Hum-Comput Stud 2012;70(11):888–906.
- [25] Bolton ML, Bass EJ. Generating erroneous human behavior from strategic knowledge in task models and evaluating its impact on system safety with model checking. IEEE Trans Syst Man Cybern Syst 2013;43(6):1314–27.
- [26] Bolton ML, Bass EJ. Evaluating human-human communication protocols with miscommunication generation and model checking. In: Proceedings of the fifth NASA formal methods symposium. moffett field: NASA ames research center. Moffett Field: NASA Ames Research Center; 2013, p. 48–62.
- [27] Pan D, Bolton ML. Properties for formally assessing the performance level of human-human collaborative procedures with miscommunications and erroneous human behavior. Int J Ind Ergon 2018;63:75–88.
- [28] Bolton ML, Bass EJ. Evaluating human-automation interaction using task analytic behavior models, strategic knowledge-based erroneous human behavior generation, and model checking. In: Proceedings of the IEEE international conference on systems man and cybernetics. Piscataway: IEEE; 2011, p. 1788–94.
- [29] Bolton ML, Molinaro KA, Houser AM. A formal method for assessing the impact of task-based erroneous human behavior on system safety. Reliab Eng Syst Saf 2019;188:168–80.
- [30] Bolton ML. Model checking human-human communication protocols using task models and miscommunication generation. J Aerosp Inf Syst 2015;12(7):476–89.
- [31] Bolton ML, Zheng X, Molinaro K, Houser A, Li M. Improving the scalability of formal human-automation interaction verification analyses that use task-analytic models. Innov Syst Softw Eng 2016;13:1–17.
- [32] Bolton ML, Bass EJ. A method for the formal verification of human interactive systems. In: Proceedings of the 53rd annual meeting of the human factors and ergonomics society. Santa Monica: HFES; 2009, p. 764–8.
- [33] Bolton ML, Bass EJ. Building a formal model of a human-interactive system: Insights into the integration of formal methods and human factors engineering. In: Proceedings of the 1st NASA formal methods symposium. Moffett Field: NASA Ames Research Center; 2009, p. 6–15.
- [34] Bolton ML, Bass EJ. Using model checking to explore checklist-guided pilot behavior. Int J Aviat Psychol 2012;22(4):343–66.
- [35] Bolton ML, Bass EJ. Using task analytic models and phenotypes of erroneous human behavior to discover system failures using model checking. In: Proceedings of the human factors and ergonomics society annual meeting. 54, (13):Los Angeles, CA: SAGE Publications Sage CA; 2010, p. 992–6.
- [36] Bolton ML. A task-based taxonomy of erroneous human behavior. Int J Hum-Comput Stud 2017;108:105–21.
- [37] Hollnagel E. The phenotype of erroneous actions. Int J Man-Mach Stud 1993;39(1):1–32.
- [38] Reason J. Human error. New York: Cambridge University Press; 1990.

- [39] Bell J, Holroyd J. Review of human reliability assessment methods. Technical report RR679, Derbyshire: Health and Safety Executive; 2009.
- [40] Di Pasquale V, Iannone R, Miranda S, Riemma S. An overview of human reliability analysis techniques in manufacturing operations. In: Schiraldi M, editor. Operations management. InTech; 2013, p. 221–40.
- [41] Swain A. Accident sequence evaluation program human reliability analysis procedure. Technical report NUREG/CR-4772, Washington, DC: US Nuclear Regulatory Commission; 1987.
- [42] Williams JC. HEART – A proposed method for achieving high reliability in process operation by means of human factors engineering technology. In: Proceedings of a symposium on the achievement of reliability in operating plant, safety and reliability society. Birmingham: NEC; 1986.
- [43] Hollnagel E. Cognitive reliability and error analysis method (CREAM). Oxford: Elsevier; 1998.
- [44] Fujita Y, Hollnagel E. Failures without errors: Quantification of context in HRA. Reliab Eng Syst Saf 2004;83(2):145–51.
- [45] Reer B. Review of advances in human reliability analysis of errors of commission part 2: EOC quantification. Reliab Eng Syst Saf 2008;93(8):1105–22.
- [46] Hollnagel E. Context, cognition and control. In: Waern Y, editor. Co-operative process management, cognition and information technology. London: Taylor & Francis; 1998, p. 27–52.
- [47] Stanton NA, Ashleigh MJ, Roberts AD, Xu F. Testing Hollnagel's contextual control model: Assessing team behaviour in a human supervisory control task. J Cogn Ergon 2001;5(1):21–33.
- [48] Blom HAP, Stroeve S, Daams J, Nijhuis HB. Human cognition performance model based evaluation of air traffic safety. In: Proceedings of the 4th international workshop on human error, safety and system development. Linköping. 2001, p. 11–2.
- [49] Worm A. Breaking the barriers: Facilitating efficient command and control in multi-service emergency management. In: 8th world conference on emergency management. Oslo. 2001, p. 19–22.
- [50] Hollnagel E, Kaarstad M, Lee H-C. Error mode prediction. Ergonomics 1999;42(11):1457–71.
- [51] Geng J, Murè S, Baldissoni G, Camuncoli G, Demichela M. Human error probability estimation in ATEX-HMI area classification: From THERP to FUZZY CREAM. Chem Eng Trans 2015;43:1243–8.
- [52] Castiglia F, Giardina M, Caravello FP. Fuzzy Fault Tree analysis in modern γ -ray industrial irradiator: Use of fuzzy version of HEART and CREAM techniques for human error evaluation. In: International conference on probabilistic safety assessment and management, 2008.
- [53] Rantanen E, Deeter J, Burke S, Wang Y. Human factors evaluation of pharmacy operations. In: 2012 Symposium on human factors and ergonomics in health care: bridging the gap. Santa Monica: HFES; 2012.
- [54] Yang Z, Bonsall S, Wall A, Wang J, Usman M. A modified CREAM to human reliability quantification in marine engineering. Ocean Eng 2013;58:293–303.
- [55] Rashed SK. The concept of human reliability assessment tool CREAM and its suitability for shipboard operations safety. J Shipp Ocean Eng 2016;6:313–20.
- [56] Chen D, Fan Y, Li W, Wang Y, Zhang S. Human reliability prediction in deep-sea sampling process of the manned submersible. Saf Sci 2019;112:1–8.
- [57] Zheng X, Bolton ML, Daly C, Feng L. A formal human reliability analysis of a community pharmacy dispensing procedure. In: Proceedings of the human factors and ergonomics society annual meeting. Los Angeles: SAGE; 2017, p. 728–32.
- [58] Zhang S, He W, Chen D, Chu J, Fan H. A dynamic human reliability assessment approach for manned submersibles using PMV-CREAM. Int J Naval Archit Ocean Eng 2019.
- [59] Zheng X, Bolton ML, Daly C, Bileteoff E. The development of a next-generation human reliability analysis: Systems analysis for formal pharmaceutical human reliability (SAFPHR). Reliab Eng Syst Saf 2020;202:15 pages. <http://dx.doi.org/10.1016/j.res.2020.106927>.
- [60] Zheng X, Bolton ML, Daly C. Extended SAFPHR (Systems Analysis for Formal Pharmaceutical Human Reliability): Two approaches based on extended CREAM and a comparative analysis. Saf Sci 2020;132. <http://dx.doi.org/10.1016/j.ssci.2020.104944>, 18 pages.
- [61] Bedford T, Bayley C, Revie M. Screening, sensitivity, and uncertainty for the CREAM method of human reliability analysis. Reliab Eng Syst Saf 2013;115:100–10.
- [62] He X, Wang Y, Shen Z, Huang X. A simplified CREAM prospective quantification process and its application. Reliab Eng Syst Saf 2008;93(2):298–306.
- [63] Di Pasquale V, Miranda S, Iannone R, Riemma S. A simulator for human error probability analysis (SHERPA). Reliab Eng Syst Saf 2015;139:17–32.
- [64] Kwiatkowska M, Norman G, Parker D. PRISM 4.0: Verification of probabilistic real-time systems. In: International conference on computer aided verification. Springer; 2011, p. 585–91.
- [65] SyncRO Soft SRL. Relax NG schema diagram. 2021, URL: <http://www.oxygenxml.com/doc/ug-oxygen/topics/relax-ng-schema-diagram.html>.
- [66] Byrne MD, Bovair S. A working memory model of a common procedural error. Cogn Sci 1997;21(1):31–61.
- [67] Ratwani RM, Gregory Trafton J. A generalized model for predicting postcompletion errors. Top Cogn Sci 2010;2(1):154–67.
- [68] Ratwani RM, Trafton JG. A real-time eye tracking system for predicting and preventing postcompletion errors. Hum-Comput Interact 2011;26(3):205–45.
- [69] Bolton ML, Bass EJ. Using task analytic models to visualize model checker counterexamples. In: Proceedings of the 2010 IEEE international conference on systems, man, and cybernetics. Piscataway: IEEE; 2010, p. 2069–74.
- [70] Abbate AJ, Bass EJ. Modeling affordance using formal methods. In: Proceedings of the human factors and ergonomics society annual meeting, vol. 61, no. 1. SAGE Publications Sage CA: Los Angeles, CA; 2017, p. 723–7.
- [71] Kim N, Shin D, Wysk R, Rothrock L. Using finite state automata (FSA) for formal modelling of affordances in human-machine cooperative manufacturing systems. Int J Prod Res 2010;48(5):1303–20.