

A FRAMEWORK TO STUDY HUMAN-AI COLLABORATIVE DESIGN SPACE EXPLORATION

Antoni Viros-i-Martin*

Daniel Selva

Systems Engineering, Architecture and Knowledge Laboratory
Department of Aerospace Engineering
Texas A&M University
College Station, Texas 77840
Email: aviros@tamu.edu

ABSTRACT

This paper presents a framework to describe and explain human-machine collaborative design focusing on Design Space Exploration (DSE), which is a popular method used in the early design of complex systems with roots in the well-known design as exploration paradigm. The human designer and a cognitive design assistant are both modeled as intelligent agents, with an internal state (e.g., motivation, cognitive workload), a knowledge state (separated in domain, design process, and problem specific knowledge), an estimated state of the world (i.e., status of the design task) and of the other agent, a hierarchy of goals (short-term and long-term, design and learning goals) and a set of long-term attributes (e.g., Kirton's Adaption-Innovation inventory style, risk aversion). The framework emphasizes the relation between design goals and learning goals in DSE, as previously highlighted in the literature (e.g., Concept-Knowledge theory, LinD model) and builds upon the theory of common ground from human-computer interaction (e.g., shared goals, plans, attention) as a building block to develop successful assistants and interactions. Recent studies in human-AI collaborative DSE are reviewed from the lens of the proposed framework, and some new research questions are identified. This framework can help advance the theory of human-AI collaborative design by helping design researchers build promising hypotheses, and design studies to test these hypotheses that consider most relevant factors.

INTRODUCTION

Human designers – be they engineers, software developers, hardware designers, industrial designers, or architects – have been aided in the design process by intelligent tools for over six decades now [1]. As systems from all domains become more complex, and given that our capabilities as human designers do not evolve at the same rate, intelligent design tools have grown in capabilities to better assist humans in the design process, leveraging advances in artificial intelligence and computing infrastructure.

In this paper, we use the term cognitive assistants (CAs) to refer to a broad class of such intelligent tools that act as agents that “augment human intellect” [2]. By casting CAs as agents, we imply that they interact with an environment and have goals they try to achieve [3]; however, this definition does not imply any level of autonomy or sophistication in the goals or in the strategies used to achieve those goals. A standard gradient-based design optimization tool can be compared to an agent with a single static goal of finding the best possible design in a static design space using a single strategy – going in the direction of the gradient; however, the definition also includes more intelligent assistants that reason (e.g., generate sub-goals) based on their internal state and estimated state of the world, use a variety of strategies to achieve goals (e.g., different heuristics), and learn (i.e., adapt to the environment to achieve their goals more efficiently). Russell and Norvig define these agents as learning agents [3].

* Address all correspondence to this author.

This paper proposes a framework that describes the interaction and collaboration between two agents – the human designer and a cognitive design assistant – involved in a design task. As such, although the framework applies to rudimentary interactions (e.g., through static Graphical User Interfaces), it is designed to capture more sophisticated interactions including interactive and natural language interfaces. These kinds of intelligent agents, which have some reasoning and learning abilities and have a natural language interface, are the focus of this paper. They are also more aligned with modern definitions of CA, such as the one by Le and Wartschinski in [4]: “Cognitive Assistants offer computational capabilities typically based on Natural Language Processing, Machine Learning, as well as reasoning chains operating on large amounts of data, enabling them to assist humans in cognitive processes”.

Although many design theories have been proposed, to the best of our knowledge there is no single cohesive theoretical framework in the literature that models the interaction and collaboration between a human designer and a CA during a design process. There are, however, models that describe parts of it in the design cognition literature [5, 6] (design by a single human individual), design teams literature [7, 8] (human-human collaboration), and human-machine collaboration literature [9, 10] (human-CA collaboration outside of design). This framework builds upon these three bodies of work and focuses on the specific context of design space exploration, as a first step towards a more general theory of human-AI collaborative design.

To create better cognitive design assistants that improve design outcomes (e.g., design quality and diversity, efficiency) we need a deeper understanding of the interaction between human designers and CAs during the design process [11]. The goal of the proposed framework is to provide a cohesive theoretical frame of reference and a common language to develop and validate new explanatory and predictive models that focus on specific components or processes of human-AI collaborative design. This is echoed by the review on design cognition by Hay et al. in [11], where they remark the importance of studying how AI can help in the design process, especially as a companion to human designers. Although the ideal framework would describe all possible interactions between human designers, intelligent tools, and design processes, this is out of scope for this paper.

A very high level view of the model is shown in Figure 1. The main components involved are the Human Designer, the Cognitive Assistant, the Tradespace Exploration Tool, and the Tradespace Exploration Problem that is being solved. We have chosen to focus on Design Space Exploration or Tradespace Exploration [12] as our design task of choice due to its popularity in the early phases of complex system design [13].

Tradespace Exploration is widely used across different types of design (e.g., consumer product design, complex system design, and architecture design), and is a useful process for two main reasons: First, by systematically and quantitatively com-

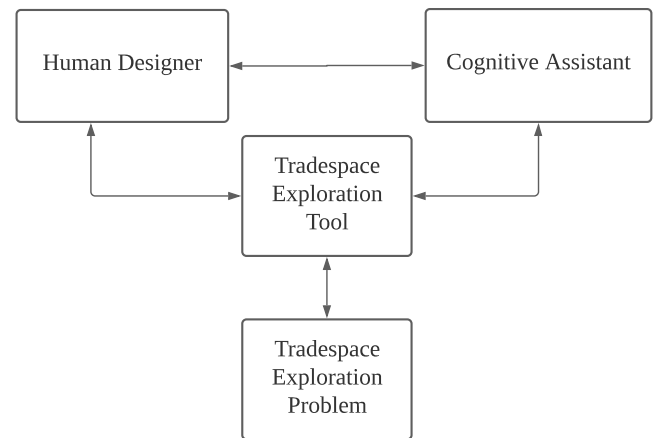


FIGURE 1. HIGH LEVEL VIEW OF THE MODEL ARCHITECTURE.

paring a large number of design alternatives, it helps reduce human biases from past experience and expertise, which can result in better design outcomes (better design performance/cost/risk obtained with less resources). Second, Tradespace Exploration can help the designer learn useful information about the design problem [6, 14] (e.g. what design decisions have the highest impacts on the final value of the designs, how changes in the needs of stakeholders affect the most valuable designs, etc.). For example, two studies with NASA’s Jet Propulsion Laboratory’s A-Team concluded that, in Tradespace Exploration studies for early space missions design, the knowledge extraction part of the process is at least as important as finding the best designs. In the first study, Fillingim et al. [15] mention how design heuristics – rules of thumb to guide the design process – are heavily used during the tradespace exploration process by professionals, to the point it justifies creating a formal repository of this knowledge so that designers can perform better in future design problems. In the second study [16], Viros-i-Martin and Selva conclude that obtaining this knowledge about the design space and the trade-offs between decisions is an important outcome of the process, as evidenced by the answers on a semi-structured interview with the test subjects (e.g., “having a clear understanding of this information is vital to our job as decisions need to be justified to humans.”) As stakeholders’ needs change during the early design process, possessing this knowledge becomes key to rapidly evolving designs to meet new needs and restrictions. Having this knowledge about the design problem can also lead to better design decisions [17, 18]. This knowledge discovery aspect of Tradespace Exploration is central to the framework and model we describe in this paper.

If we go into a more detailed view of the process, the objec-

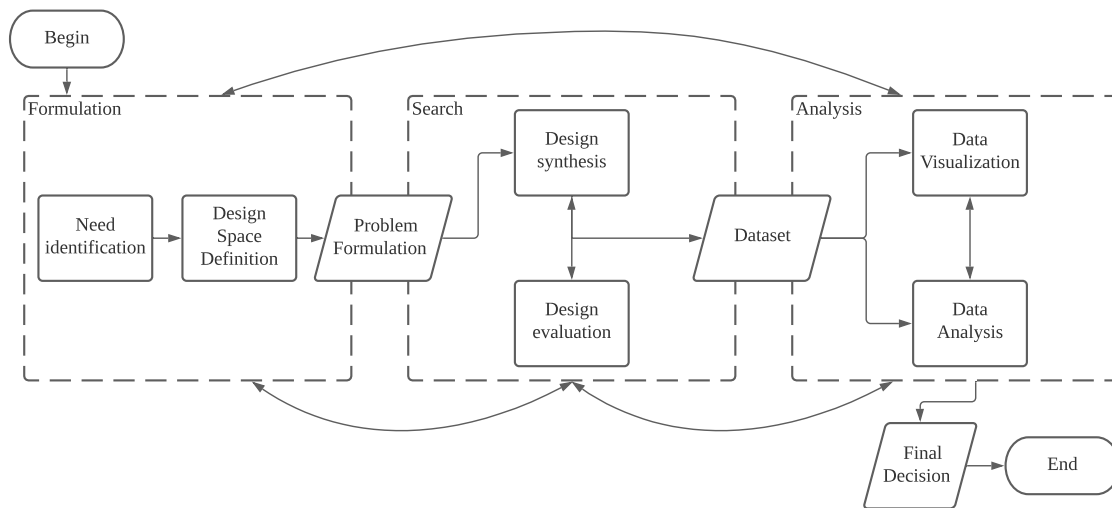


FIGURE 2. TRADESPACE EXPLORATION DESIGN PROCESS.

tive of Tradespace Exploration as a design process is to choose one or a set of design alternatives from a pool of possibly millions. Each design alternative is described as a set of decisions that can be organized differently for each design problem: decisions can be discrete or continuous, and sets of decisions might be grouped as assignments, combinations, partitions, etc. To better describe what we mean by design decisions and their grouping, we follow the example of a satellite system design problem. Nonetheless, this design process can be applied to any design problem that can be decomposed into decisions. To make the example more clear, as well as show all the steps and paths involved in the process, we introduce Figure 2. Following our example, first we start on the Formulation step, by identifying the needs of stakeholders in the satellite system (e.g. they want to create a satellite system that measures soil moisture at the Earth's surface by using well-tested components). This is followed by the decision of how many satellites will the system have, as well the decision on what structure will the satellites follow in space (constellation, formation, train, federation, etc). This results in a Problem Formulation, with which the Search step can begin. Then, the next decision is what to put in each satellite and where should that satellite be, which amounts to design synthesis. Some of these decisions are discrete, while others are continuous, and some of the decisions can be organized in sets that translate to assignment or partition problems. Once that is performed, each constellation is evaluated for cost and metrics related to stakeholders' needs, and this process is repeated until enough of them exist, by which point we have a Dataset. A design dataset can be visualized and analyzed, which are the two main activities in the Analysis step. For example, one might look for patterns in

well-performing satellite systems. Based on the results of this analysis, the problem can either be reformulated by going back to the first step, or new designs can be synthesized based on the insights obtained during the process. If the designer is content with the results on both designs and learning, the process can end with final design choices; which in our example would amount to the satellite system that will end up going to space. This zig-zag process in tradespace exploration is common in design processes based on the "design as exploration" philosophy that will be detailed later in the paper [19–21].

In the framework we are describing in this paper, the concept of knowledge and the acquisition of it – learning – take the foremost role. A unified representation of knowledge can help when creating robust metrics for measuring its acquisition. One example of this pair of knowledge representation and a set of metrics for measuring it can be observed in [22], where Bang and Selva present Knowledge Graphs (KG) as an explicit representation for knowledge and create a set of metrics and tests to measure learning. In our framework, we do not prescribe any knowledge representation for both human and CA's knowledge, but we do require that one exists, that it is explicit, and that metrics exist for learning to be measured with that representation of knowledge.

The presented framework is of explanatory rather than predictive nature. It is intended to help design researchers explain the phenomena observed in studies in the literature. It also describes certain kinds of interactions that have never been observed yet due in part to lack of capabilities of current cognitive design assistants. For most of the framework description, we will have a running example based on Daphne [23], a design

CA developed to help space engineers design distributed satellite missions for Earth Observation.

The rest of the paper is organized as follows: First, we go through the literature related to the main concepts and building blocks of our framework. Following that, we describe the main aspects of the framework by going into each part of it and detailing the chosen models and attributes for both the human designer, the CA, and their relationships. Then, we discuss how the framework can describe the tradespace exploration process with a temporal view of the interactions between the agents involved, including examples of different kinds of interactions. Finally, we discuss the implications of this model in the design of future CAs, as well as plans for its validation.

RELATED WORK

Our framework builds on theories from different fields. In this section, we review works on design cognition, design theory, and human-machine interaction – from robots to computers to CAs, in settings as different as a warehouse and a concurrent design facility. We also review part of the large body of work on understanding designer teams, with a strong focus on interpersonal communication and how that affects each step of the design process.

Design Theories and Design Cognition

Design as a science is usually traced back to the seminal book from Simon, “The Sciences of the Artificial” [24], published five decades ago. The design theory in the book describes a design paradigm based on understanding the design activity as a problem-solving activity. This translates to considering “design as search”: when given a problem statement, one has to search for an optimal or satisficing solution. This breakthrough motivated many design studies to test various aspects related to the theory. Dinar et al. [25] describe over 180 related studies published since the publication of Simon’s book. In the last 25 years, new theories of design have appeared, but most if not all of them describe one or more “zig-zagging” processes, where the designer jumps back and forth between two related processes or spaces. For example, design as co-evolution [19] describes a zig-zagging process between problem space and solution space, and Concept-Knowledge (C-K) theory [20] describes a zig-zagging process between the concept space and the knowledge space. The Function-Behaviour-Structure theory describes design as zig-zagging between the design function and its structure, through different processes involving the behavior of the design [21]. All of these processes sit on top of the zig-zagging process between synthesis and evaluation of designs implicit in any iterative design search scheme. Collectively, all these zig-zagging processes and their interplay are what constitute the “design as exploration” framework, implemented in various related

processes such as Trade Space Exploration [12] and set-based design [26].

In the recent past, three different review papers [25, 27, 28] have come to the conclusion that to design better tools for designers understanding the design process is not enough. What is needed is a better understanding of the cognitive processes in the human designer’s mind. The main roadblock to improving the understanding of a designer’s mind, according to the reviews, is the lack of a common vocabulary or ontology for the design processes, as well as the fact that metrics about cognition are difficult to measure. On the one hand, an attempt at unifying the diverse design cognition theories in the literature is a theory by Cash and Kreye known as the Uncertainty Driven Action (UDA) model [29], where design is viewed as a process with three actions, all related to uncertainty perception. The driver for the human to go through the design process is to reduce their uncertainty with relation to the design problem, be it by acquiring information, sharing it with other agents involved in the design, or creating a representation for the design. On the other hand, a new field of studying designers’ brains directly is opening up, as described by Gero and Milovanovic in [30].

Cognitive Assistants

The structure most modern CAs follow can be traced back to a DARPA project known as the Cognitive Assistant that Learns and Organizes (CALO) [31]. Many commercial CAs follow the structure defined in CALO for CAs. Siri, from Apple, is a direct spin-off of the project. As described in the Introduction, modern CAs – as well as CALO – offer computational capabilities typically based on Natural Language Processing, Machine Learning, and reasoning chains operating on large amounts of data, with which they assist humans [4]. Of note, CAs do not need to implement a cognitive architecture such as Soar [32] or ACT-R [33] – and most of them currently do not implement any –, and neither do they need to learn from their interactions with users in any form. This being said, our vision of CAs for this model is more of a design peer that teams with the human designer rather than a chatbot with no reasoning capabilities, which requires some human understanding capabilities. This means at least listening and adapting to the designer to align priorities and goals during the design process. The most general model that contains our vision of a CA is that of Intelligent Agents [3], which observe their environment, and then act on it based off goals they are trying to achieve. To formally model the attributes of the CA, we look at the taxonomy for CAs defined in [34] by Maier et al. In that paper, they define CAs in terms of 4 main characteristics: Learning, Intelligence, Autonomy, and Communication. Each of these 4 characteristics has a scale unique to them, with Gagne’s Hierarchy of Learning [35] for Learning, Bloom’s Taxonomy [36] for intelligence, an adapted version of the 5-Level Classification Scale for Autonomous Vehicles [37] for Autonomy, and the Hi-

erarchy of Natural Language Processing Skills [38] for Communication. For example, Daphne, the CA the authors have worked on, has a rating of “Concept Learning” in the Learning attribute – due to its ability to consistently answer queries with the same intent even if the phrasing is different each time –, “Evaluate” in the Intelligence attribute – due to its ability to criticize designs from the user –, “Subsystem” in the Autonomy attribute – due to its inability to completely carry out the design task without human intervention –, and “Multiple Meanings” in the Communication attribute – due to its usage of advanced Speech-To-Text techniques –. In our framework, we further specialize some of these scales to make them more relevant to the design problem.

Usage of CAs in design – as per the modern definition provided above – is a pretty new development, so not a lot of them are available publicly. All the design CAs mentioned in this paragraph inform our model, both because of their strengths and their weaknesses. The list includes the Daphne family of CAs, with a version for Earth Observation mission design [23], and a version for analysis of simulation data of Martian Entry, Descent, and Landing [39], as well as the Systems Engineering Advisor [40], whose purpose is to find gaps in requirements for space missions. Other CAs include the intelligent agents in Hyform [41], geared towards acting as peers to a hybrid design team for tasks involving Unmanned Aerial Vehicles, the intelligent virtual agent [10], which helps with designing electoral districts, and the Design Engineering Assistant [42], geared towards data processing to help with early space mission design. There are examples of design assistants that do not fit exactly with the modern definition of a CA yet are still relevant to our model as examples of certain interactions. The Architect Collaborator (TAC) [43], conducts a tradespace exploration of building designs for architects based on requests from them. The Intelligent Manufacturing System Design Assistant (IMDSA) [44] uses expert systems to assist in the design of manufacturing plants. Sindi [45] performs tradespace analysis based on designer preferences for highways. PQE and Q-Chef [46] are a framework and implementation of an assistant that finds novel designs based on the user input for recipes. All of these CAs are mentioned during the paper when their capabilities are a relevant example to the framework.

Human-Machine Collaboration

In our framework, we are trying to model the interaction between a human and a machine. The field with the most research in the area, although with little attention to the design task, is robotics. One of the key processes identified in human-robot interaction as key to communication is the process of accumulating common knowledge, beliefs, and assumptions in a collaborative task, known as *common ground* [47]. Common ground helps collaborators know what information their partners need, how to present information so that it is understood, and whether partners have interpreted information correctly [47]. The field of

robotics has studied common ground as the basis for collaboration between humans and robots [48]. Hoffman and Breazeal [9] describe human-robot collaboration as based on shared goals and plans, where the robot must be able to adapt in real-time to the human actions. They claim that, based on the Joint Intention Theory [49], in order for effective collaboration to emerge, all teammates must have shared beliefs about the state of the task, and a coordinated plan of action, as well as trust in the counterpart. A successful application of this theory can be identified in [50], where Nikolaidis et al. create a reinforcement learning model based on joint action observation that significantly improves team performance. In recent years, the research in the field has moved towards fluency in the interaction, which is a measure of how synchronized these hybrid teams of humans and machines are [51]. Although design has not been the focus of human-robot interaction, some recent work has looked into it. For example, Cobbie [52] is a CA turned robot to foster creativity in sketches, and in [53], Law et al. present a robotic arm that does tradespace exploration together with a human. Finally, design research has shown that adverse effects exist when creating AIs for design that do not take into account the human characteristics and intentions [54].

To use all this corpus of robotics research in the field of design, our framework maps the concepts of shared goals and plans between human and machine to the task of design space exploration, listing the goals that human and machine move towards, as well as how the CA plans according to its prediction of what the human designer is planning.

Design Teams

Design can happen in teams. Our framework describes a hybrid team composed of a human and a CA – extensions to multiple humans and one or more CAs are left for future work. A lot of research has gone into understanding how design teams work to make them better at performing the design task, although usually those teams are composed solely of multiple humans. Our framework’s task is to adapt the results of the many works on design teams to a hybrid team of human and CA. This research has been approached from many fronts: for example, on the topic of team formation, Jablowski et al. [8] investigate the effects of team composition in terms of cognitive characteristics such as Kirton’s Adaption-Innovation (KAI) scale, which measures an individual’s cognitive preference for structure in generating and working with ideas in problem solving. They find that teams with individuals with a wide range of KAI scores discuss and explore a higher variety of concepts, and they also find that encouraging certain behaviors can increase the number of unique ideas generated. Also related to KAI, KABOOM [55] is an agent-based framework that simulates design teams with different individual values of KAI scores and communication rates to see how effective they are in a design task without requiring humans in the

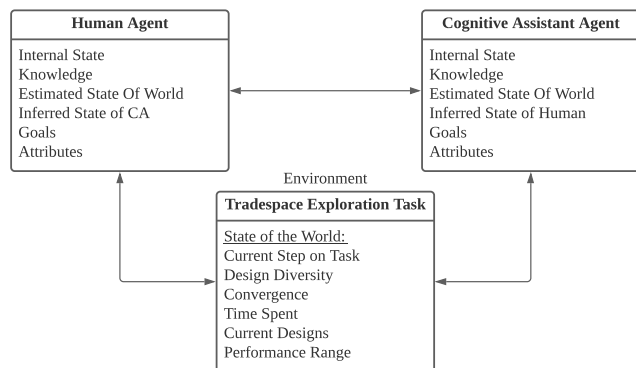


FIGURE 3. OVERVIEW OF THE FRAMEWORK: THE HUMAN DESIGNER AND THE CA ARE MODELED AS INTELLIGENT AGENTS.

loop. In Teamology [56], Wilde explains how using the Myers-Briggs personality types can help in creating more effective design teams. Yet another approach is based on measuring the personality of team members and how that changes team dynamics using the “Big Five Factors”, as seen in [57–60]. Our framework uses these metrics to better describe the human side of the team, which in turn helps the CA adapt.

On the topic of team dynamics, we focus on the concept of “roles”. In [7], Paton and Dorst describe that highly experienced design teams tends to assign different roles to everyone on the team, with the most common ones being “technician”, “facilitator”, “expert/artist” and “collaborator”. Similar (almost identical) roles also came up during our study of NASA’s JPL A-Team [16]. The concept of “roles” is key to our description of a design CA, as we have found CAs in the literature play one or more of the roles described here during the design space exploration task, usually changing between two or three of them when helping the human designer.

MODEL OF HUMAN-AI COLLABORATIVE DSE

The pieces that compose the framework are shown in Figure 3, a more detailed view of Figure 1. The main idea is that both the Human Designer and the Cognitive Assistant are modelled as Intelligent Agents [61], while the Tradespace Exploration Task is modelled as the environment where both agents exist and interact.

The remainder of this section will be devoted to describing the characteristics of both agents in Figure 3, describing Tradespace Exploration as an environment, and showing an example of both agents’ behavior during the design task. To better illustrate the concepts introduced in this section and the next, we will use an example based on using the Daphne CA [23] to

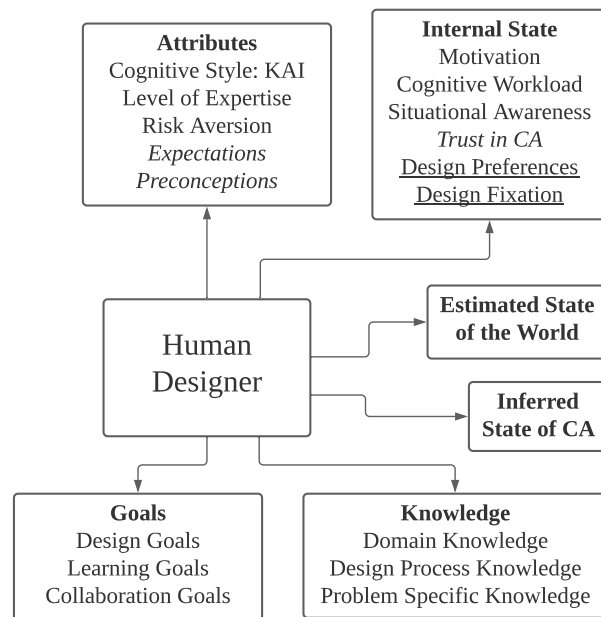


FIGURE 4. DESCRIPTION OF THE HUMAN DESIGNER.

explore the design space for a next-generation space mission to measure soil moisture. Occasionally, we will switch to another of the many CAs for design mentioned in the literature review as needed to illustrate specific points.

The Human Designer

The human designer is modeled as an intelligent agent [61]. A standard model for intelligent agents includes three layers, whose outputs are fed to the next layer as input: perception, decision, and action. In this framework, the perception layer’s input is both the state of the design task – state of the world – as well as the CA’s interactions, while its output is the human’s estimated state of the world, as well as an inferred internal state of the CA. With this information, together with its own internal State, Attributes, Goals, and Knowledge about the field and design task (all described below), the human agent makes a decision on what to do next. Finally, the designer acts on the task or interacts with the CA based on what they have decided to do. All of the human’s internal characteristics are summarized in Figure 4: Attributes, which are characteristics of the agent that do not change during a single design task; Internal State, which is a set of characteristics that can and will evolve during the design process; Estimated State of the World, which is an approximation of the state of the design task; Inferred State of the CA, which is what the human thinks the CA’s state looks like, given their interactions; Knowledge, of which we focus on Domain Knowledge,

Design Process Knowledge, and Problem Specific Knowledge; and Goals, which guide the agent actions during the design process.

Many of the cognitive theories for design mentioned in the literature review fit the intelligent agent model: all of them have the concept of Goals and Actions as key parts of the human cognition. Combining this with the Tradespace Exploration Task environment and the human characteristics summarized in the Attributes and Internal State makes for a complete intelligent agent model. Modern CAs – or at least those based on the architecture of CALO, which includes most of them – are designed with the assumption that the humans interacting with them act in accordance to the Belief-Desire-Intention (BDI) [62] cognitive model for intelligent agents. Finally, this choice of modeling also allows for parallels with the CA, which is modeled as an intelligent agent as well.

The choice of attributes for the human designer is based on what the literature on design cognition and human-machine collaboration has found to be important when working on tradespace exploration tasks, as well as the interaction with the CA in the design process, as that is what the model is built upon. In Figure 4, characteristics marked in *italics* have a strong relation to the interaction with the CA, and mostly come from the Human Machine Collaboration field, while those underlined are closely related to the design task and come primarily from the design theory literature. The following list enumerates the attributes we have chosen for our model and why we have chosen them: 1) *Cognitive Style (KAI)*, because prior research has found that teams with diverse scores in KAI tend to increase exploration of the tradespace [8]. Therefore, measuring a designer's KAI score can guide the CA's actions; 2) *Level of Expertise*, because studies reviewed in [25] show that the behavior of design experts is very different to that of novices. Incidentally, this has also been observed in studies performed with cognitive design assistants [16, 63]; 3) *Risk Aversion*, because a designer's bias towards certain design choices can be influenced by the human's aversion to innovative but unproven designs, as shown in many studies reviewed in [25]; 4) *Expectations and Preconceptions*, i.e., how a human expects a machine to act based both on what they know of similar machines and their past experiences, which according to the HCI literature can significantly affect the collaboration [9, 64].

The choice of internal state for the human designer is, once again, based on what the literature has found to be important for design space exploration and collaboration with machines. An issue about all these short-term metrics is that measuring them in real-time is hard, as most quantitative metrics require the usage of surveys. These are the variables in the designer's state: 1) *Motivation*, because prior literature mentions that emotions, and especially motivation, should be a part of the design process [30]; 2) *Cognitive Workload*, because of findings in HCI literature that put a cap on the amount of cognitive load before performance

starts to degrade in intensive tasks [65]; 3) *Situational Awareness*, because it is a relevant metric for complex, rapidly changing environments such as tradespace exploration according to the literature [66]. Additionally, the CA can guide the human designer back to a highly aware state if they become disoriented; 4) *Trust in automation*, because from HCI literature and some results with cognitive design assistants [63], we know it correlates to performance in the design task; 5) *Design Preferences*, defined here as both the human designer's preferences on how to perform exploration of the design space as well as their preferred design decisions, whether based on facts or guts, because several of the reviewed protocol studies in [25] highlight their importance in relation to the design outcomes. They can change during the design process, based on new knowledge obtained about a design; 6) *Design Fixation*, again because several studies in design cognition, as mentioned in [25], have found it to be correlated with design task performance. We must mention that Design Fixation is usually understood as a phenomenon, so we are abusing the language to denote the degree of Design Fixation of the designer during the design process.

One specially important part of the human designer's state in our model is their Knowledge, which justifies why it is a separate box in Figure 4. For the purpose of our framework, and based on the knowledge classification in [67], we focus on the following parts of the designer's knowledge: Problem-Specific knowledge, Domain knowledge, and Design Process knowledge. Problem Specific knowledge is, as its name suggests, specific to the current design task at hand – e.g. about what good designs look like, features driving the structure of the design space, or sensitivities of criteria to decisions, for that problem instance. Domain knowledge concerns all general knowledge related to a specific field – e.g. space and spacecraft, aerodynamics and planes, thermodynamics and engines – which may be learned during the DSE task. Design Process knowledge refers to general knowledge about the design process – in our case, Design Space Exploration. For example, how to interpret sensitivity indices, or why a Pareto Front might lose usability when there are too many objectives.

Finally, concerning goals, we distinguish between Design goals and Learning goals as observed in Figure 5, with a hierarchy based on the time horizon of that goal, as described in [6]. The most important goals in this framework are the ones in the middle of the hierarchy, which have the design project as their time horizon. On one hand, the goal of any design task is to find a design that maximizes value. On the other hand, there is the goal of learning as much as possible about the specific design problem that is being solved. In the long term, the organization for which the designer is working for or the designer themselves might also have goals, which can be as varied as creating long-lasting design ideas or generalizing design-specific knowledge into field knowledge that can be reused in future designs. A proposition of this model is that the relationship between design and learning goals

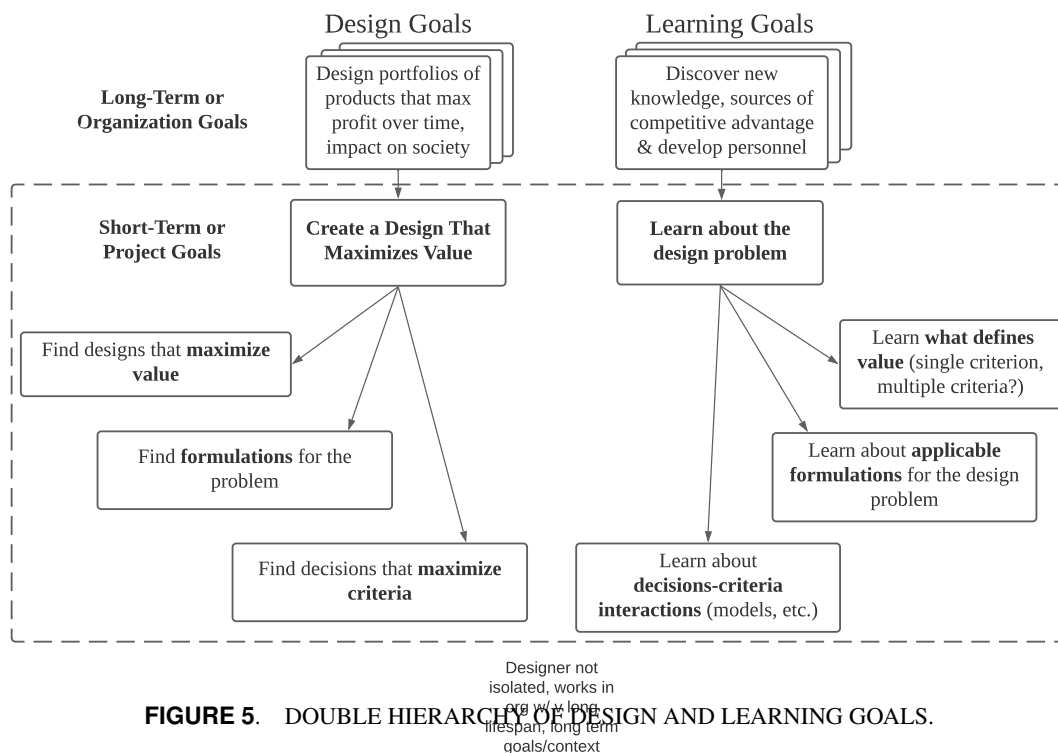


FIGURE 5. DOUBLE HIERARCHY OF DESIGN AND LEARNING GOALS.

is synergistic. Both task-level goals can also be decomposed in shorter-term goals, all of which further the main design or learning goal. The human designer will always be focused on trying to fulfill one or more of them. For example, a designer trying to design a new satellite system for Earth Observation first learns the value of the system based on what their mission requirements are and which stakeholders are involved, then looks for valid representations of the system so that new systems can be designed that maximize the value that has been defined. Although this example will be expanded upon in the example timeline, we can see how the human designer switches between design and learning goals continuously. This goal structure ties well with both the UDA [5] and LinD [6] models of design cognition, as well as the different steps and products generated during Tradespace Exploration, which will be described in its own subsection. Apart from these hierarchies, both agents have Collaboration goals, which inform how the approach the interaction with the other agent. In the case of the human designer, these goals are informed by their Expectations and Preconceptions with respect to the CA agent.

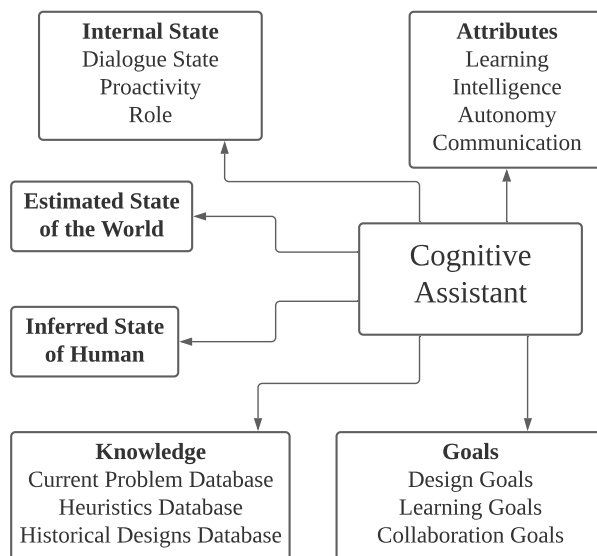


FIGURE 6. COGNITIVE ASSISTANT STRUCTURE.

The Cognitive Assistant

In our framework, the CA is also modeled as an intelligent agent, and therefore it also has the perception, decision, and action layers from the intelligent agent model. Similarly, as observed in Figure 6, its characteristics can be decomposed in its Attributes, its Internal State, its Estimated State of the World,

its Inferred State of the Human Designer, its Knowledge, and its Goals.

The attributes are taken from [34], as it is a relatively complete taxonomy of cognitive assistants, including scales for each

of the attributes with explanations for each level of each attribute. Learning refers to the degree to which a CA can connect ideas to create new knowledge, ranging from simply responding to queries to full-fledged Problem Solving for unique problems. Intelligence refers to the degree a CA can process information, and ranges from remembering facts to creation of new information. Autonomy refers to the ability of the CA to act without waiting for explicit permission from the human (e.g., a request), and can range from completely passive to fully autonomous. Finally, Communication refers to its Natural Language Processing capabilities, ranging from labelling words or sentences to being able to understand analogies. Daphne, as our example, was already classified in the Related Work section.

The internal state for the CA depends on each CA's implementation. For our framework, we list the following basic components: 1) A Dialogue State that can vary from simply storing the last question or query from the user to a full-blown dialogue history, together with an up-to-date context of the conversation to be able to converse with more agility and accuracy. Some CAs do not have conversation capabilities, but an Action History can be used instead or at the same time; 2) the Proactivity Level can either be fixed or a slider from never taking initiative without being queried to completely guiding the design process instead of the human designer; 3) the Role, inspired by human designer teams and their specialized roles, is the current specialized behavior the CA is exhibiting, and different roles that are commonly implemented in published CAs will be described shortly. We assume a CA can only play a single role at a time. Daphne, as it stands right now, is able to keep a limited context of the conversation (e.g., what pronouns such as "this" may refer to), has a fixed proactivity level for each design session for a given role, and has 5 different roles.

The framework organizes the knowledge state of the CA in 3 different databases. The Current Problem Database contains all information related to the current problem: there is a representation of the design set for the tradespace exploration task, including the problem formulation and all the designs discovered under that formulation, as well as all the knowledge that comes from mining the dataset, such as the driving features for certain regions, value sensitivities to decisions, etc. The Heuristics Database is a collection of expert knowledge related to the design task at hand, and is a combination of Domain Knowledge and Design Process Knowledge as described above. Finally, the Historical Designs Database is a collection of past results of both the current and past design tasks that can be searched through for insights related to the current design, such as what decisions have been taken before. Our example CA, Daphne, has the three databases.

The CA's Design and Learning goals follow the same structure as the human's – see the last subsection for a more complete definition. This ensures that shared goals and plans are possible between both agents. Apart from those, the CA also has Collab-

oration Goals, which are related to helping the human designer through the design task. Examples of this goal include reducing the Design Fixation of the human – by suggesting unexplored regions of the design space, for example –, increasing the common ground to ensure a fluid collaboration, or increasing the human trust in the CA. Daphne currently does not have its own explicit goals or goal-setting capabilities, but its implicit goals are to improve the design set and guide the human designer through critique of their designs.

It is worth to get into a detailed description of roles, a part of the internal state of design CAs. The following is a non-exhaustive list of roles that current CAs in the literature play during the design process.

"Historian" role: When a design CA plays the role of Historian, it looks up past designs in design databases to find insights that are useful to the human designer. Such insights might be heuristics, what heuristics are useful to the current problem vs which are not, or even if a design has been attempted before. These insights can be presented on-demand in response to a human query or pro-actively to guide the design process. CAs with this role: Daphne-EO [23], ESA's DEA [42].

"Analyst" role: When a design CA plays the role of Analyst, it performs data analysis on the current tradespace, trying to find patterns in regions of interest and presenting them to the human designer. These patterns can be found using any of the multitude of data mining or feature extraction algorithms that currently exist. These data-driven insights can be shown to the user by request or pro-actively. CAs with this role: Daphne-EO, Daphne-EDL [39], HyForm [41].

"Explorer" role: When a design CA plays the role of Explorer, it helps the human designer find new, relevant designs that have not been found yet. The methods to do this are varied, and include global and local search and optimization strategies. Again, these new designs can be shown to the user by request or pro-actively. CAs with this role: Daphne-EO, HyForm, Law's IVA [10], TAC [43], [45].

"Expert" role: When a design CA plays the role of Expert, it answers questions about the value of designs and the models used to calculate it, such as why the CA is providing a certain score for a given design. As with all other roles, it can be either reactive or proactive. CAs with this role: Daphne-EO, HyForm, SEA [40], IMDSA [44].

"Critic" role: When a design CA plays the role of Critic, it criticizes existing designs, identifying strengths and weaknesses, and giving suggestions on how to improve it. The Critic role is particular in that its function is to gather and aggregate input from the other roles. For example, the Expert in Daphne may provide the Critic with a weakness of a design that is violating a rule of thumb about designing Earth observation satellites, while the Historian may point

out to the Critic that a design appears to be very novel compared to the historical design database. Through this aggregation function, the Critic role realizes something closer to a real peer designer. As always, it can work re-actively or pro-actively. CAs with this role: Daphne-EO, PQE [46].

Tradespace Exploration as an Environment

In Figure 3, the Tradespace Exploration Task is modelled as the environment where both the human and CA agents interact. As such, it has a set of attributes that both agents will observe to create their estimated versions of the state of the world. If we start by looking at Figure 2, one of these attributes is the current step in the design process: the agents can either be in Formulation, Search, Analysis, or picking the Final Choices. Continuing with the same Figure, both the Problem Formulation and the Dataset are also part of the state of the world. The last important metric, specially given that tradespace exploration is usually performed with tight time requirements and that it can affect the behavior of the agents, is the time spent on the task – or the time left to finish it. To complete the state of the world, we include other attributes such as performance metrics (e.g. performance range, cost range) and design diversity metrics (e.g. convergence, crowding distance).

Example Timeline

Figure 7 showcases examples of interactions between the human designer and the cognitive assistant during different steps of the design process of a satellite system. The CA represented in the example is an ideal CA that implements all functions and capabilities described in the framework. Currently, only a CA simulated through a Wizard-of-Oz experience [68] could provide all the interactions shown in the example. Theoretically, any of the cognitive architectures we mentioned previously – such as Soar and ACT-R –, as well as more general architectures such as Reinforcement Learning agents can be trained to perform all the capabilities we mentioned. The main challenge remains in the difficulty and resources needed to train the agents. The point of this example is to showcase a wide range of interactions, show the interlinks between different zig-zagging processes in tradespace exploration, and highlight the relation between learning and designing goals.

On the range of interactions, we see the CA taking on the five roles we described in the framework. For example, in the first interaction, it acts as an Expert, and then switches to being an Explorer to find new designs. Later, it plays the roles of Analyst, Historian, and Critic. The CA can rapidly change its role in an interaction by interaction basis. These examples also showcase different levels of proactivity, from a small degree at the beginning by asking follow-up questions to a completely autonomous exploration of the tradespace, back to a reactive behavior when asked about trends.

If we focus on the steps of the design process that are described in the middle of the Figure, we can observe how, even though the process starts with Formulation and ends with Final Choices, the process itself has the “zig-zag” behavior we have previously mentioned, with jumps between the steps of Formulation, Search, and Analysis.

In the right timeline, we can observe the human switching goals between Learning and Design, and up and down the hierarchy described in Figure 5. As already described, both agents’ goals for the task at hand are learning as much as possible about the design problem while also creating the most valuable designs. What we observe on the timeline is how the human’s goal at each moment of time has a narrower scope than the ultimate goal. For example, in the first interaction, the human is preoccupied with finding out the requirements for the design they are doing, while the assistant helps by giving out those requirements as well as extra questions to learn about the applicable representations. Later, the human designer wants to mine the data for insights to create designs that maximize the value to the stakeholders. These insights compel the human to make changes to the problem formulation. Finally, at the end of the design process, the human designer wants to justify their final design choice to stakeholders, going back to the top of the hierarchy for both design and learning goals. On the CA side of things, we can observe the agent tracks the human designer’s goals, and tries to keep the interactions fluid as part of its collaboration goals.

When talking about learning, we can also look at it from the lens of a model such as LinD [6]. LinD describes a taxonomy of design actions with their relation to learning. In the example in Figure 7, we can see examples of information gathering actions – such as the first interaction –, synthesising – when the CA critiques the human’s design –, or analysing – when looking for rules that describe the behavior of design in a certain region of the tradespace –. LinD further classifies most of these actions as in-situ learning activities, given that learning occurs at the time of the design action, not before or after.

CONCLUSION

In this paper, we have presented a framework that describes the interactions between human designers and cognitive assistants that occur during the tradespace exploration process in the early design of complex systems. We have described both the human designer and the cognitive assistant as intelligent agents and related their characteristics to previous research into human cognition in design, human-machine collaboration, and design teams, as well as described their interactions during the tradespace exploration process through an example.

The next steps for this research in the short-term are both enhancing the behavioral part of the framework by formalizing it, and performing a more exhaustive search of the literature to create a complete list of important attributes for both the human and



model of the human inside the CA that looks similar to the one described in this paper.

This would allow for a systematic validation of the framework, with features and attributes being enabled or disabled for different control groups to ascertain whether each of the components mentioned in the model are important and consistent with theory. Results from these studies may result in modifications to the framework, e.g., if we realize an important aspect is missing.

In the longer term, this framework would ideally evolve from an explanatory model for hypothesis-building to a predictive model that can anticipate what types of interactions are likely to result in better design and learning outcomes. Armed with those models, we hope we can design truly useful and insightful design CAs that realize the vision of a peer designer.

Ultimately, the objective of this framework is to aid in the creation of CA agents that improve the design process, including, for example, an increase of design diversity, and improvement of design convergence, and any other metric that might be deemed crucial for a particular design problem.

ACKNOWLEDGMENT

We would like to acknowledge funding from NSF CMMI grant number 1907541.

REFERENCES

- [1] Hayes, C. C., Goel, A. K., Tumer, I. Y., Agogino, A. M., and Regli, W. C., 2011. "Intelligent Support for Product Design: Looking Backward, Looking Forward". *Journal of Computing and Information Science in Engineering*, **11**(021007), June.
- [2] Engelbart, D. C., 1962. Augmenting human intellect: a conceptual framework. Tech. rep., Stanford Research Institute, Washington D.C.
- [3] Russell, S., and Norvig, P., 2020. *Artificial Intelligence: A Modern Approach*, 4th edition ed. Pearson, Hoboken, Apr.
- [4] Le, N.-T., and Wartschinski, L., 2018. "A Cognitive Assistant for improving human reasoning skills". *International Journal of Human-Computer Studies*, **117**, Sept., pp. 45–54.
- [5] Cash, P. J., 2018. "Developing theory-driven design research". *Design Studies*, **56**, May, pp. 84–119.
- [6] Sim, S. K., and Duffy, A. H. B., 2004. "Evolving a model of learning in design". *Research in Engineering Design*, **15**(1), Mar., pp. 40–61.
- [7] Paton, B., and Dorst, K., 2011. "Briefing and reframing: A situated practice". *Design Studies*, **32**(6), Nov., pp. 573–587.
- [8] Jablow, K. W., Sonalkar, N., Edelman, J., Mabogunje, A., and Leifer, L., 2019. "Investigating the Influence of Designers' Cognitive Characteristics and Interaction Behaviors in Design Concept Generation". *Journal of Mechanical Design*, **141**(091101), Apr.
- [9] Hoffman, G., and Breazeal, C., 2004. "Collaboration in Human-Robot Teams". In AIAA 1st Intelligent Systems Technical Conference, Infotech@Aerospace Conferences, American Institute of Aeronautics and Astronautics.
- [10] Law, M. V., Kwatra, A., Dhawan, N., Einhorn, M., Rajesh, A., and Hoffman, G., 2020. "Design Intention Inference for Virtual Co-Design Agents". In Proceedings of the 20th ACM International Conference on Intelligent Virtual Agents, IVA '20, Association for Computing Machinery, pp. 1–8.
- [11] Hay, L., Cash, P., and McKilligan, S., 2020. "The future of design cognition analysis". *Design Science*, **6**. Publisher: Cambridge University Press.
- [12] Ross, A. M., and Hastings, D. E., 2005. "The Tradespace Exploration Paradigm". Accepted: 2014-01-27T15:52:50Z.
- [13] Crawley, E., Cameron, B., and Selva, D., 2015. *System Architecture: Strategy and Product Development for Complex Systems*, 1st edition ed. Pearson, Boston, Apr.
- [14] Grecu, D. L., and Brown, D. C., 1996. "Dimensions of learning in agent-based design".
- [15] Fillingim, K. B., Nwaeri, R. O., Borja, F., Fu, K., and Paredis, C. J. J., 2020. "Design Heuristics: Extraction and Classification Methods With Jet Propulsion Laboratory's Architecture Team". *Journal of Mechanical Design*, **142**(081101), Feb.
- [16] Viros i Martin, A., and Selva, D., 2020. "Daphne: A Virtual Assistant for Designing Earth Observation Distributed Spacecraft Missions". *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, **13**, pp. 30–48. Conference Name: IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing.
- [17] Simpson, T. W., Carlsen, D., Malone, M., and Kollat, J., 2011. "Trade Space Exploration: Assessing the Benefits of Putting Designers "Back-in-the-Loop" during Engineering Optimization". In *Human-in-the-Loop Simulations: Methods and Practice*, L. Rothrock and S. Narayanan, eds. Springer, London, pp. 131–152.
- [18] Zhang, X. L., Simpson, T., Frecker, M., and Lesieutre, G., 2012. "Supporting knowledge exploration and discovery in multi-dimensional data with interactive multiscale visualisation". *Journal of Engineering Design*, **23**(1), Jan., pp. 23–47. Publisher: Taylor & Francis eprint: <https://doi.org/10.1080/09544828.2010.487260>.
- [19] Maher, M. L., Poon, J., and Boulanger, S., 1996. "Formalising Design Exploration as Co-Evolution". In *Advances in Formal Design Methods for CAD: Proceedings of the IFIP WG5.2 Workshop on Formal Design Methods for Computer-Aided Design, June 1995*, J. S. Gero and F. Sudweeks, eds., IFIP — The International Federation for Information Processing. Springer US, Boston, MA, pp. 3–30.
- [20] Hatchuel, A., and Weil, B., 2003. "A new approach of innovative Design: an introduction to CK theory".
- [21] Gero, J. S., and Kannengiesser, U., 2004. "The situated function-behaviour-structure framework". *Design Studies*, **25**(4), July, pp. 373–391.
- [22] Bang, H., and Selva, D., 2020. "Measuring Human Learning in Design Space Exploration to Assess Effectiveness of

- Knowledge Discovery Tools". American Society of Mechanical Engineers Digital Collection.
- [23] Bang, H., Viros, A., Prat, A., and Selva, D., 2018. "Daphne: An Intelligent Assistant for Architecting Earth Observing Satellite Systems".
- [24] Simon, H. A., 1996. *The sciences of the artificial (3rd ed.)*. MIT Press, Cambridge, MA, USA.
- [25] Dinar, M., Shah, J. J., Cagan, J., Leifer, L., Linsey, J., Smith, S. M., and Hernandez, N. V., 2015. "Empirical Studies of Designer Thinking: Past, Present, and Future". *Journal of Mechanical Design*, **137**(021101), Feb.
- [26] Singer, D. J., Doerry, N., and Buckley, M. E., 2009. "What is set-based design?". *Naval Engineers Journal*, **121**(4), pp. 31–43. Publisher: American Society of Naval Engineers.
- [27] Hay, L., Duffy, A. H. B., McTeague, C., Pidgeon, L. M., Vuletic, T., and Greal, M., 2017. "A systematic review of protocol studies on conceptual design cognition: Design as search and exploration". *Design Science*, **3**. Publisher: Cambridge University Press.
- [28] Cash, P., Hicks, B., and Culley, S., 2015. "Activity Theory as a means for multi-scale analysis of the engineering design process: A protocol study of design in practice". *Design Studies*, **38**, May, pp. 1–32. Publisher: Elsevier.
- [29] Cash, P., and Kreye, M., 2017. "Uncertainty Driven Action (UDA) model: A foundation for unifying perspectives on design activity". *Design Science*, **3**. Publisher: Cambridge University Press.
- [30] Gero, J. S., and Milovanovic, J., 2020. "A framework for studying design thinking through measuring designers' minds, bodies and brains". *Design Science*, **6**. Publisher: Cambridge University Press.
- [31] Myers, K., Berry, P., Blythe, J., Conley, K., Gervasio, M., McGuinness, D. L., Morley, D., Pfeffer, A., Pollack, M., and Tambe, M., 2007. "An Intelligent Personal Assistant for Task and Time Management". *AI Magazine*, **28**(2), June, pp. 47–47. Number: 2.
- [32] Laird, J. E., III, R. E. W., Wang, Y., Derbinsky, N., Nuxoll, A. M., Lathrop, S., Wintermute, S., III, R. P. M., Gorski, N., and Xu, J., 2012. *The Soar Cognitive Architecture*. The MIT Press, Cambridge, Mass. ; London, England, Apr.
- [33] Ritter, F. E., Tehranchi, F., and Oury, J. D., 2019. "ACT-R: A cognitive architecture for modeling cognition". *WIREs Cognitive Science*, **10**(3), p. e1488. eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1002/wcs.1488>.
- [34] Maier, T., Menold, J., and McComb, C., 2019. "Towards an Ontology of Cognitive Assistants". *Proceedings of the Design Society: International Conference on Engineering Design*, **1**(1), July, pp. 2637–2646. Publisher: Cambridge University Press.
- [35] Soulsby, D., 1975. "Gagne's Hierarchical Theory of Learning: Some Conceptual Difficulties". *Journal of Curriculum Studies*, **7**(2), Nov., pp. 122–132. Publisher: Routledge. eprint: <https://doi.org/10.1080/0022027750070204>.
- [36] Anderson, L., Krathwohl, D., Airasian, P., Cruikshank, K., Mayer, R., Pintrich, P., Rath, J., and Wittrock, M., 2000. *Taxonomy for Learning, Teaching, and Assessing, A: A Revision of Bloom's Taxonomy of Educational Objectives, Abridged Edition*, 1st edition ed. Pearson, New York, Dec.
- [37] J3016B: Taxonomy and Definitions for Terms Related to Driving Automation Systems for On-Road Motor Vehicles - SAE International.
- [38] Richard, G. J., 2016. *The Source for Processing Disorders*, 2nd edition ed. Pro Ed, Austin, Texas, Oct.
- [39] León, S. S. D., Selva, D., and Way, D. W., 2019. "A Cognitive Assistant for Entry, Descent, and Landing Architecture Analysis". In 2019 IEEE Aerospace Conference, pp. 1–12. ISSN: 1095-323X.
- [40] Salado, A., 2020. "From Model-Based Requirements to a Virtual Systems Engineering Advisor that Identifies Gaps in Requirements: An Application to Space Systems". In ASCEND 2020, American Institute of Aeronautics and Astronautics. eprint: <https://arc.aiaa.org/doi/pdf/10.2514/6.2020-4092>.
- [41] Song, B., Zurita, N. F. S., Zhang, G., Stump, G., Balon, C., Miller, S. W., Yukish, M., Cagan, J., and McComb, C., 2020. "Toward Hybrid Teams: A Platform to Understand Human-Computer Collaboration During the Design of Complex Engineered Systems". *Proceedings of the Design Society: DESIGN Conference*, **1**, May, pp. 1551–1560. Publisher: Cambridge University Press.
- [42] Berquand, A., Murdaca, F., Riccardi, A., Soares, T., Generé, S., Brauer, N., and Kumar, K., 2019. "Artificial Intelligence for the Early Design Phases of Space Missions". In 2019 IEEE Aerospace Conference, pp. 1–20. ISSN: 1095-323X.
- [43] Koile, K., 2004. "An Intelligent Assistant for Conceptual Design". In Design Computing and Cognition '04, J. S. Gero, ed., Springer Netherlands, pp. 3–22.
- [44] Floss, P., and Talavage, J., 1990. "A knowledge-based design assistant for intelligent manufacturing systems". *Journal of Manufacturing Systems*, **9**(2), Jan., pp. 87–102.
- [45] Mandow, L., and Pérez-de-la Cruz, J.-L., 2004. "Sindi: an intelligent assistant for highway design". *Expert Systems with Applications*, **27**(4), Nov., pp. 635–644.
- [46] Grace, K., Maher, M. L., Wilson, D., and Najjar, N., 2017. "Personalised Specific Curiosity for Computational Design Systems". In Design Computing and Cognition '16, J. S. Gero, ed., Springer International Publishing, pp. 593–610.
- [47] Clark, H. H., and Brennan, S. E., 1991. "Grounding in communication". In *Perspectives on socially shared cognition*. American Psychological Association, Washington, DC, US, pp. 127–149.
- [48] Stubbs, K., Hinds, P. J., and Wettergreen, D., 2007. "Au-

- tonomy and Common Ground in Human-Robot Interaction: A Field Study”. *IEEE Intelligent Systems*, 22(2), Mar., pp. 42–50. Conference Name: IEEE Intelligent Systems.
- [49] Cohen, P. R., and Levesque, H. J., 1991. “Teamwork”. *Noûs*, 25(4), pp. 487–512. Publisher: Wiley.
- [50] Nikolaidis, S., Ramakrishnan, R., Gu, K., and Shah, J., 2015. “Efficient Model Learning from Joint-Action Demonstrations for Human-Robot Collaborative Tasks”. In 2015 10th ACM/IEEE International Conference on Human-Robot Interaction (HRI), pp. 189–196. ISSN: 2167-2121.
- [51] Hoffman, G., 2019. “Evaluating Fluency in Human–Robot Collaboration”. *IEEE Transactions on Human-Machine Systems*, 49(3), June, pp. 209–218. Conference Name: IEEE Transactions on Human-Machine Systems.
- [52] Lin, Y., Guo, J., Chen, Y., Yao, C., and Ying, F., 2020. “It Is Your Turn: Collaborative Ideation With a Co-Creative Robot through Sketch”. In Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems, CHI ’20, Association for Computing Machinery, pp. 1–14.
- [53] Law, M. V., Jeong, J., Kwatra, A., Jung, M. F., and Hoffman, G., 2019. “Negotiating the Creative Space in Human-Robot Collaborative Design”. In Proceedings of the 2019 on Designing Interactive Systems Conference, DIS ’19, Association for Computing Machinery, pp. 645–657.
- [54] Zhang, G., Raina, A., Cagan, J., and McComb, C., 2021. “A cautionary tale about the impact of AI on human design teams”. *Design Studies*, 72, Jan., p. 100990.
- [55] McComb, C., Jablow, K., and Lapp, S., 2019. KA-BOOM: An Agent-Based Model for Simulating Cognitive Sale in Team Problem Solving. Tech. rep., engrXiv, July. type: article.
- [56] Wilde, D. J., 2009. *Teamology: The Construction and Organization of Effective Teams*. Springer-Verlag, London.
- [57] Kichuk, S. L., and Wiesner, W. H., 1997. “The big five personality factors and team performance: implications for selecting successful product design teams”. *Journal of Engineering and Technology Management*, 14(3), Sept., pp. 195–221.
- [58] Ayşar, A. Z., Valencia-Romero, A., and Grogan, P. T., 2019. “The Effects of Locus of Control and Big Five Personality Traits on Collaborative Engineering Design Tasks With Negotiation”. American Society of Mechanical Engineers Digital Collection.
- [59] Toh, C. A., Patel, A. H., Strohmetz, A. A., and Miller, S. R., 2016. “My Idea Is Best! Ownership Bias and its Influence on Engineering Concept Selection”. American Society of Mechanical Engineers Digital Collection.
- [60] Stidham, H., Flynn, M., Summers, J. D., and Shuffler, M., 2018. “Understanding Team Personality Evolution in Student Engineering Design Teams Using the Five Factor Model”. American Society of Mechanical Engineers Digital Collection.
- [61] Franklin, S., and Graesser, A., 1997. “Is It an agent, or just a program?: A taxonomy for autonomous agents”. In Intelligent Agents III Agent Theories, Architectures, and Languages, J. P. Müller, M. J. Wooldridge, and N. R. Jennings, eds., Lecture Notes in Computer Science, Springer, pp. 21–35.
- [62] Georgeff, M., Pell, B., Pollack, M., Tambe, M., and Wooldridge, M., 1999. “The Belief-Desire-Intention Model of Agency”. In Intelligent Agents V: Agents Theories, Architectures, and Languages, J. P. Müller, A. S. Rao, and M. P. Singh, eds., Lecture Notes in Computer Science, Springer, pp. 1–10.
- [63] Viros i Martin, A., and Selva, D., 2020. “Learning Comes from Experience: The Effects on Human Learning and Performance of a Virtual Assistant for Design Space Exploration”.
- [64] Kwon, M., Jung, M. F., and Knepper, R. A., 2016. “Human expectations of social robots”. In 2016 11th ACM/IEEE International Conference on Human-Robot Interaction (HRI), pp. 463–464. ISSN: 2167-2148.
- [65] Wickens, C., and Tsang, P. S., 2015. “Workload”. In *APA handbook of human systems integration*, APA handbooks in psychology®. American Psychological Association, Washington, DC, US, pp. 277–292.
- [66] O’Brien, K. S., and O’Hare, D., 2007. “Situational awareness ability and cognitive skills training in a complex real-world task”. *Ergonomics*, 50(7), July, pp. 1064–1091. Publisher: Taylor & Francis. eprint: <https://doi.org/10.1080/00140130701276640>.
- [67] Selva Valero, D., 2012. “Rule-based system architecting of Earth observation satellite systems”. Thesis, Massachusetts Institute of Technology. Accepted: 2013-01-07T21:19:40Z.
- [68] Dahlbäck, N., Jönsson, A., and Ahrenberg, L., 1993. “Wizard of Oz studies — why and how”. *Knowledge-Based Systems*, 6(4), Dec., pp. 258–266.