

Research Article

Cite this article: Bhattacharya K, Majumder A, Bhatt AN, Keshwani S, BSC R, Venkataraman S and Chakrabarti A (2024). Developing a method for creating structured representations of working of systems from natural language descriptions using the SAPPhIRE model of causality. *Artificial Intelligence for Engineering Design, Analysis and Manufacturing*, **38**, e24, 1–13
<https://doi.org/10.1017/S0890060424000118>

Received: 01 June 2023

Revised: 03 May 2024

Accepted: 12 June 2024

Keywords:

creative design; causal ontology; natural language processing; design by analogy; analogical reasoning; SAPPhIRE model

Corresponding author:

Kausik Bhattacharya;

Email: kausikb@iisc.ac.in

Developing a method for creating structured representations of working of systems from natural language descriptions using the SAPPhIRE model of causality

Kausik Bhattacharya¹ , Anubhab Majumder¹, Apoorv Naresh Bhatt¹ ,
Sonal Keshwani², Ranjan BSC³, Srinivasan Venkataraman⁴  and
Amaresh Chakrabarti¹

¹Department of Design and Manufacturing, Indian Institute of Science, Bangalore, India; ²Human Centered Design, Indraprastha Institute of Information Technology, Delhi Okhla Phase III, New Delhi; ³Department of Mechanical Engineering, Dr. H. N. National College of Engineering, Bangalore, India and ⁴Department of Design, Indian Institute of Technology Delhi (IITD), New Delhi, India

Abstract

Due to their significant role in creative design ideation, databases of causal ontology-based models for biological and technical systems have been developed. However, creating structured database entries through system models using a causal ontology requires the time and effort of experts. Researchers have worked toward developing methods that can automatically generate representations of systems from documents using causal ontologies by leveraging machine learning (ML) techniques. However, these methods use limited, hand-annotated data for building the ML models and have manual touchpoints that are not documented. While opportunities exist to improve the accuracy of these ML models, more importantly, it is required to understand the complete process of generating structured representations using causal ontology. This research proposes a new method and a set of rules to extract information relevant to the constructs of the SAPPhIRE model of causality from descriptions of technical systems in natural language and report the performance of this process. This process aims to understand the information in the context of the entire description. The method starts by identifying the system interactions involving material, energy and information and then builds the causal description of each system interaction using the SAPPhIRE ontology. This method was developed iteratively, verifying the improvements through user trials in every cycle. The user trials of this new method and rules with specialists and novice users of the SAPPhIRE modeling showed that the method helps in accurately and consistently extracting the information relevant to the constructs of the SAPPhIRE model from a given natural language description.

1. Introduction

Design by Analogy is a powerful method to support creativity in product design (Nagai and Taura, 2015). Therefore, many databases based on system models, such as DANE or IDEA-INSPIRE, were developed to support design by analogy in product design (Fu et al., 2014). However, creating accurate system models using an ontology like SBF (Goel et al., 2009) or SAPPhIRE (Chakrabarti et al., 2005) requires considerable time and effort from specialists. Based on the literature, the automation of the process for creating the SBF or SAPPhIRE models using the information in technical documents is attempted without reporting the complete process of how to build these models accurately. Therefore, there exists scope for understanding this modeling process, as this knowledge will be useful in accurately producing structured representations and can help in choosing an appropriate automation strategy.

This research aims to develop a basis for support development that can automatically generate causal, ontology-based representations from descriptions of systems in natural language. Therefore, the scope of this research is on the process that must precede its automation. The goals of this research are: (a) to understand the process for creating representation using the SAPPhIRE model of causality from descriptions in natural language and (b) to develop a new method based on this understanding to extract information about the entities and relationships among them of the SAPPhIRE model from natural language descriptions of systems. This research will help select a suitable automation strategy in the future.

The paper starts with a detailed literature review in Section 2. Since this research belongs to the broader field of Design Creativity and Design by Analogy, Section 2.1 describes the prior research in these two areas, highlighting the significance of the choice of ‘representation’ for design stimulus. Section 2.2 describes the efforts to support Design by Analogy. Since a design stimulus

© Kausik Bhattacharya, 2024. Published by Cambridge University Press. This is an Open Access article, distributed under the terms of the Creative Commons Attribution licence (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted re-use, distribution and reproduction, provided the original article is properly cited.

can be represented in several ways, multiple types of design support containing structured representations of stimuli have been developed. As a causal representation of systems' functioning described using the SAPPhIRE model is the scope of this research, [Section 2.3](#) explains the SAPPhIRE model. The goal of this research is to lay the groundwork for a support that can automatically convert descriptions in natural language into causal ontology-based representations. So, [Section 2.4](#) explains the prior art in the automatic generation of structured representations, including causal representations. The research questions and the research methodology are described in [Section 3](#) and the research findings are presented in [Sections 4](#) and [5](#). [Section 6](#) presents the conclusions and the future scope of this research.

2. Literature review

2.1 Design creativity and design-by-analogy

A design is a means for changing an existing situation into a preferred one (Simon, 1969). Creativity in design is often characterized as a process by which an agent (like a designer) uses its ability to generate outcomes that are novel and useful (Chakrabarti, 2009). Researchers studied the influence of different design methods on creative design outcomes (Chulvi et al., 2012) and developed methods for enhancing creativity. Stimuli are important in innovative product design (Howard et al., 2010). Literature provides evidence that using an external stimulus compared to no stimulus can generate more ideas (Srinivasan et al., 2017). Systematic use of knowledge from artificial and natural domains helps designers generate many solutions and develop them into realizable and practical prototypes (Sarkar et al., 2008). A detailed study to understand designers' approach to select inspirational stimuli in the conceptual design phase was conducted where the representation of external stimuli was an essential factor in the inspiration process (Gonçalves et al., 2016). "Design-by-analogy or Analogical Design is the process of developing solutions through mapping of attributes, relations, and purposes that a source problem or situation may share (or at least partially share) with an existing target solution or situation" (Goel, 1997). Analogical reasoning is considered the core of human cognitive activities responsible for creativity in design (Goel and Shu, 2015). Many empirical studies were conducted to understand the process and the factors influencing the effectiveness of Design-by-Analogy (Srinivasan et al., 2015; Helms et al., 2009; Song et al., 2019; Keshwani et al., 2015; Keshwani, 2018). Representation of a stimulus plays an important role in Design-by-Analogy (Hashemi Farzaneh et al., 2015; Lenau et al., 2015) and it can influence fixation and creativity during idea generation (Moreno et al., 2016; Bonnardel, 2000; Vandevienne et al., 2016; Verhaegen et al., 2011).

2.2 Support for design-by-analogy

Search for analogues is a key step in the Design-by-Analogy process (Sartori et al., 2010) because finding relevant analogues can influence the nature of design outcomes. This paper outlines some of the commonly cited Design-by-Analogy support. Support like the Functional Model database (Nagel et al., 2010) and AskNature (Deldin et al., 2014) use a function-based approach to identify analogues from a database. The Functional Model database has an engineering-to-biology thesaurus that maps biological terms to the functional basis of technical systems. AskNature categorizes the information of the biological functions according to the four layers of Biomimicry Taxonomy (Group, Sub-group, Function, and

Strategy). Biological models are mapped into different engineering fields using "strategy". Analogy Retriever (Han et al., 2018) uses 16 ontological relationships to describe the connections between various system entities. This supports analogical reasoning for creative ideation by solving proportional analogy problems.

On the other hand, supports like Functional Vector (Murphy et al., 2014) and SEABIRD (Vandevienne et al., 2016b) use the vector space method to find analogues based on the semantic similarity of words. In the Functional Vector model, a query vector of functions is generated. A relevancy score of the query with a functional vocabulary is calculated using a patent database. SEABIRD method generates the Product Aspects (PA) and Organism Aspects (OA) matrices from a database of technical system documents and biological functions. Mapping between two domains is then quantified based on the values from the mathematical product of the PA and OA matrices.

A third category of support, like DANE (Vattam et al., 2011) and IDEA-INSPIRE (Chakrabarti et al., 2005), uses ontology-based data models. DANE uses data queries at multiple levels of abstractions in a controlled database comprising structured data models of SBF (i.e., Structure-Behavior-Function) (Goel et al., 2009). In SBF, 'Structures' are the constituent components, substances, and relations among them; 'Behavior' is the series of state changes from an input to an output state, and the transition from one state to another happens through functions. 'Function', therefore, is used as a behavioral abstraction. In IDEA-INSPIRE, the search strategy uses single or multiple levels of abstraction of the SAPPhIRE (State change – Action – Parts – Phenomenon – Input – oRgan – Effect) model (Chakrabarti et al., 2005), see [Section 2.3](#).

2.3 SAPPhIRE model of causality

The SAPPhIRE model has seven layers of abstraction, namely, State Changes, Actions, Parts, Phenomena, Inputs, oRgans, and Effects (Chakrabarti et al., 2005). They together cover the system's physical components, interfaces, interactions, structural context, and scientific laws governing them and can produce a rich comprehensive description of biological and technical systems. The SAPPhIRE model with an example (heat transfer from a hot body to cool surrounding air) is shown in [Figure 1](#). The model is scalable to represent a complex system through a multi-instance-linked model (Siddharth et al., 2018). The SAPPhIRE model can be used for many design tasks; the model can be used in design synthesis and analysis (Chakrabarti, 2009), a database of SAPPhIRE models can be used as a stimulus for Bio-inspired Design (Fu et al., 2014) and assessing novelty (Srinivasan et al., 2010; Jagtap, 2018).

2.4 Structured representation from a natural language description

Finding many stimuli and maintaining diversity and variety of content are important requirements (Srinivasan et al., 2010; Yargin et al., 2015). Databases with ontology-based models were developed to support the design creativity (Fu et al., 2014). Though ontology-based models are effective, they are handcrafted and limited in number. On the other hand, there is an abundant source of knowledge of biological and technical systems available in the form of technical documents and many websites where information is not associated with any specific format or data model. Since structured data is preferred in Design-by-Analogy because of computational advantages (Salvatore et al., 2003), researchers proposed methods using the numerical techniques of ML and data sciences that can generate structured data or

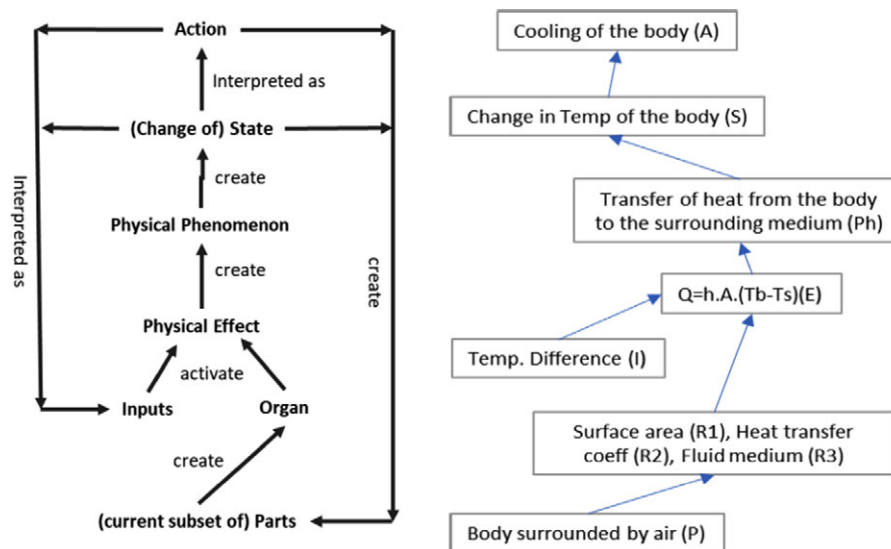


Figure 1. SAPPhIRE model of causality with an example of heat transfer from a hot body to cool surrounding air (Chakrabarti et al., 2005; Chakrabarti, 2009).

knowledge representation from information given in natural language descriptions (Jiang et al., 2022). Such methods are essentially information extraction methods, which involve skimming through documents and looking for occurrences of a particular class of information and the relationships among them (Russel & Norvig, 2010). Our ultimate motive is to determine the process for information extraction methods that can help extract information relevant to the SAPPhIRE model of causality accurately and repeatedly before it is automated. Hence, in this paper, we investigate the commonly cited approaches for information extraction and understand the underlying complete process. The references given in this section, therefore are not an exhaustive list but a representation of the dominant information extraction methods that researchers are pursuing.

A natural language processing (NLP)-based technique was developed to find causal relations between biological functions using a linguistic pattern of biologically meaningful keywords (Cheong et al., 2012). It progressively refines search keywords until a suitable match is found in a simple template to capture causal relations of biological functions (Shu, 2010). A co-occurrence graph is another powerful information extraction method, which is used to identify relevant technical parameters for a certain domain by comparing thesauri automatically extracted from patents (Cascini et al., 2011). Semantic webs are created by extracting concepts, instances, and their relationships from a handbook of the engineering domain (Hsieh et al., 2011). Rule-based techniques are also used to generate an engineering knowledge graph from a patent database with greater size and coverage (Siddharth et al., 2022). Semantic networks are considered effective for knowledge representation from vast data sources, so many such networks were developed (Chen et al. 2020; Han et al., 2022). Though the large technology semantic knowledge graph, TechNet, is reported to be better than other semantic networks such as WordNet or ConceptNet for engineering applications, more study is needed with TechNet or similar networks (Han et al., 2022; Sarica & Luo, 2021; Sarica et al., 2023).

Ontology-based models can provide multiple levels of abstractions of a system and can be used to explain systematically the functioning of a system (Srinivasan et al., 2012). Domain ontology can be extracted using design themes (things that are of the

designer's interest) from a set of technical documents (Li and Ramani, 2007). Therefore, ontology-based information extraction (OBIE) is another type of information extraction (Wimalasuriya and Dou, 2010). A two-step process using a supervised learning technique is developed where the first step identifies the occurrences of the candidate semantic entities and then identifies the ontology-specific relationships between them in the next step (Nédellec et al., 2009). Text mining techniques can be used to define a domain ontology itself, which is subsequently used in knowledge extraction and visualization (Yang et al., 2018). The biologically inspired adaptive growth approach (BIAG) learns domain ontologies from engineering documents utilizing the similarity in domain ontology learning process with a tree growth (Zhang et al., 2020). Due to its extensive knowledge representation, the ConceptNet semantic network with 16 generic ontological relations is also used to create an ontology-based representation (Han et al., 2018). However, how to use ConceptNet semantic network or knowledge graph for classical design ontologies such as SBF (Goel et al., 2009) or SAPPhIRE (Chakrabarti et al., 2005) has not been reported.

Causal representations of systems play a significant role in creative design ideation (Baldussu et al., 2012). Hence, this research focuses on an ontology that supports causal reasoning in design, a central theme of Analogical Reasoning (Lee and Holyoak, 2008; Waldmann and Hagmayer, 2013). Results show that causal models, such as SAPPhIRE or SBF are effective in design ideation (Sarkar et al., 2008; Kim et al., 2011). Hence, we pay attention to methods that automatically create representations of text using a causal ontology, such as SAPPhIRE or SBF. A support called IBID (Intelligent Bio-Inspired Design), was developed to convert natural language description into SBF-based description (Goel et al., 2020). IBID uses a combination of knowledge and machine learning-based approaches to extract and represent the knowledge using the SBF model. This work on IBID reported a detailed comparative study of multiple ML algorithms used to classify knowledge using Structure, Behavior and Function tags. Another research reported a four-step process for converting natural language descriptions into descriptions structured with the SAPPhIRE model (Keshwani et al., 2017; Keshwani, 2018). This work extracts all the sentences with potential SAPPhIRE constructs first. After that, they are split into words or collections of words. In this work, the researchers developed a

classifier using the support vector machine of the ML technique to classify the SAPPPhIRE label of each word or clause.

These last two forms of support were developed to create the SBF and SAPPPhIRE ontology-based descriptions from natural language text, which are closely aligned with the research goal of this paper. These supports use supervised learning with hand-annotated data for training and validation. However, none of these explains how to create hand-annotated data accurately. Moreover, these supports are semi-automated with major decision-making touchpoints that are currently executed manually by specialists. These manual touchpoints are not documented. The entire process of creating a structured data model from a natural language description and its overall accuracy with design examples is not reported. Therefore, the complete process and the main factors influencing the creation of ontology-based models need to be adequately researched. In an empirical study comparing three approaches, namely keyword search using Ask Nature (handcrafted database), NLP-based approach (unstructured corpus) and one where Biologists manually perform the search in literature, it is found that though the NLP-based technique is promising, it is currently prone to errors (Willox et al., 2020). This implies that it is important to understand the process before applying the NLP techniques.

3. Research question and research methodology

3.1. Research question

The aim is to find the process that can serve as the basis for developing automation for creating SAPPPhIRE models from information in natural language documents. Hence, the commonly used information extraction methods reported in the literature are studied. We find that 'ML' based methods were developed that worked along with manual steps for creating causal representations from systems descriptions in natural language. Research opportunities exist in understanding and documenting this entire process so that this understanding can act as the basis for accurately creating causal models before automating the process. Hence, this research investigates (1) understanding the entire process of extracting information and representing this information in the desired model and (2) developing an improved process. Since the SAPPPhIRE model of causality can be used to provide a rich description of biological and technical systems for design analysis and synthesis, we used it to develop and test the process. This improved manual process will be automated only after it becomes accurate and repeatable.

Therefore, the research question in this research is:

"How to create an accurate and repeatable SAPPPhIRE ontology-based representation to describe the working of technical systems from a natural language description?"

This main research question is further divided into the following two sub-questions:

1. What is the current process of extracting information and developing SAPPPhIRE representation from a given natural language description about the working of a system?
2. How should the process be modified for improving information extraction and developing SAPPPhIRE representation from a natural language description?

3.2. Research methodology

This research adopts an iterative and user-centric approach as shown in Figure 2, based on the design research methodology

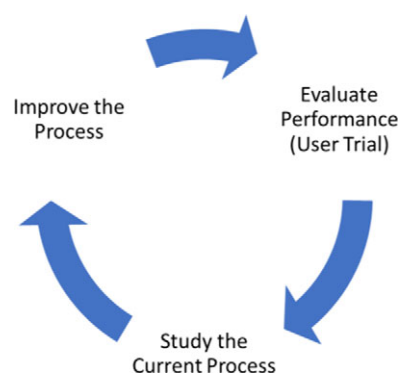


Figure 2. Research cycle.

framework of Blessing and Chakrabarti (2009). The cycle starts with understanding and improving the current process by addressing the gaps. Each research cycle produces a revised process by working with the users of the SAPPPhIRE model. The performance of the revised process is then evaluated through user trials. This cycle is repeated multiple times until a process that delivers the target performance is reached.

Since there are two research sub-questions involved, the results are presented in two parts. The first part summarizes the research carried out with the SAPPPhIRE specialists (Bhattacharya et al., 2023). The first part aims to understand the current process of generating the SAPPPhIRE representation. In the first part, improvement opportunities are identified and a process is proposed inspired by the cognition of the SAPPPhIRE specialists. The second part of the research develops the above process further by generalizing it. Hence, the second research cycle works with novice users of the SAPPPhIRE model. The reasoning rules for SAPPPhIRE information extraction are developed and validated as part of this research. These rules will be useful for later development of support for automation.

4. Developing an understanding of the current process

4.1. Intercode reliability study with SAPPPhIRE specialists

The first step is to understand the current process of developing a SAPPPhIRE model from a natural language description of the working of a system. The work for this step is carried out by working with the SAPPPhIRE specialists. These specialists are researchers who have used the SAPPPhIRE modeling for at least four years in their research. This work starts with an observational study by asking the four SAPPPhIRE specialists to generate the SAPPPhIRE representations from natural language descriptions of four different systems. The four specialists produce the SAPPPhIRE-based representations based on their understanding of the SAPPPhIRE modeling (Chakrabarti et al., 2005) and interpretation of the descriptions in natural language. The data from these models are then used to calculate and report the intercode reliability study. The procedure of this observational study is given in Appendix A. This study uses the natural language descriptions of four systems, namely, 'Elephant Turbinate', 'Bombardier Beetle', 'Thermal Wheel', and 'Electric Horn' and represented by the sample ID 'EG1', 'EG2', 'EG3' and 'EG4' respectively. The contents of these samples, given in Appendix A, are hand-curated with information taken from commonly available websites such as [howstuffworks.com](https://www.howstuffworks.com), [asknature.org](https://www.asknature.org), and Wikipedia.

Table 1. Intercooder reliability study as per current SAPPhIRE modeling

Sample ID	Total no. of words	A minimum of 3 out of 4 specialists agreed		A minimum of 2 out of 4 specialists agreed	
		Count	%	Count	%
EG1	71	6	8%	21	30%
EG2	69	2	3%	8	12%
EG3	68	1	1%	6	9%
EG4	221	2	1%	22	10%
Overall	429	11	3%	57	13%

Table 1 shows the results of the intercoder reliability of this study. The results in Table 1 show the differences among the specialists regarding the SAPPhIRE representation of the system descriptions. The first column shows the Sample ID, as explained above. For example., The sample ID EG1 represents the case of how ‘Elephant Turbine’ functions. Results for each Sample ID, shown in individual rows, are described in the subsequent columns. The second column in each row gives the total number of words relevant for all the seven SAPPhIRE constructs in a given sample, as identified by the four specialists. For example, for EG1, all four specialists identified 71 words (given in the second column) relevant to the SAPPhIRE constructs. The subsequent columns show the number of words and their percentage (%) out of the total number of words given in the second column, where a minimum of three and two out of four specialists had agreed. For example, for EG1 out of the 71 words, in 6 of them (i.e., 8%), at least 3 out of 4 specialists had an agreement (given in columns 3 and 4), and in 21 of them (i.e., 30%), at least 2 out of 4 specialists had an agreement (given in columns 5 and 6).

Workshops are then conducted to understand the reasons behind the differences. In the absence of any defined process for SAPPhIRE modeling, interpretational differences played a role in interpreting the content of the natural language text among the users. Interpretation can happen at an individual word or sentence level, at multiple sentence level, or the entire document level. The text in natural language often describes the technical process at a high level and does not provide enough details. However, the current process does not assess what information is given and what is missing in the document. It was found that due to differences in the interpretation of the text, the definition of the SAPPhIRE constructs is also not applied consistently. It is also observed that different technical or scientific terminologies are used, based on the specialist’s background knowledge, to represent the same physical behavior.

The details of this user study, including the intercoder reliability analysis, are reported by Bhattacharya et al. (2023).

4.2. Identifying improvement opportunities for the current process

Workshops with SAPPhIRE specialists after the first intercoder reliability study give inputs on how the consistency across the models could be improved. Inspired by the cognition of the SAPPhIRE specialists, a process called New Process (Preliminary) is proposed to extract SAPPhIRE-related information from technical documents in natural language. Another observational study is

then carried out with the same SAPPhIRE specialists and the New Process (Preliminary) to check if the models’ consistency improved. In this process improvement, the first goal was to understand the information given in the context of the entire document before extracting individual information relevant to each SAPPhIRE construct. The process proposed above, therefore, has two parts:

- Part A: This part is about Knowledge Understanding and Semantic Extraction. Here, the objective is to understand the information given in the context of the whole document, not just the sentence. A knowledge graph-based approach was recommended for this part, keeping in mind future automation. An example of a Knowledge Graph representing information relevant to SAPPhIRE from a system description document is given in Appendix B. However, generating the knowledge graph is not mandatory while manually executing the process. While manually performing this task, the agent (designer in this case) must understand the information in the context of the whole document instead of the context of the sentence at hand.
- Part B: This part is about extracting SAPPhIRE-related information. Here, rule-based reasoning is performed by the agent to identify information that is relevant to each SAPPhIRE construct. For automation in the future, these rules will be converted into a propositional logic of an expert system.

User trials of this new process with specialists, carried out using the procedure given in Appendix A, show that this process has the potential to improve the consistency of SAPPhIRE information extracted. Intercoder reliability analysis with specialists before and after process development is given in Table 2. Description of the columns given in the explanation of Table 1 hold goods for Table 2 as well. Table 2 shows that the score (% of SAPPhIRE words/clauses where 3 out of 4 specialists agreed) improved from 3% to 61%. Analysis of the results gives the improvement opportunities that need to be looked into before generalizing the process. The sample IDs ‘EG1’, ‘EG2’, ‘EG3’ and ‘EG4’ in Table 2 are the same as in Table 1. Sample IDs ID ‘EG5’, ‘EG6’, ‘EG7’ and ‘EG8’ are new samples representing descriptions of ‘Electric Battery’, ‘Solar Heater’, ‘Mechanical Lock’ and ‘Visualizing Infrared Rays by Fish’, respectively and given in Appendix A. The details of the process improvement opportunities based on the workshop with SAPPhIRE specialists and the user trials conducted with the new process are reported in (Bhattacharya et al., 2023). This process developed based on the workshop with the SAPPhIRE specialists is called the New Process (Preliminary) since it needs to be improved further and is not generalized yet.

It should be noted that the intercoder reliability score implies the level of repeatability in the SAPPhIRE model. A low intercoder reliability score does not necessarily imply low accuracy. The accuracy of these SAPPhIRE models was not checked individually since each of these was developed by a SAPPhIRE specialist, which was taken as a marker for accuracy.

5. Developing a new SAPPhIRE modeling procedure

5.1. Generalizing the process

For the second research sub-question, the New Process (Preliminary) described in Section 4.2 is taken as the starting point and developed further to make it a ‘to-be’ process for the future. As the work progresses on developing the process and generalizing it, the following items are addressed:

Table 2. Intercooder reliability scores WITHOUT and WITH the new process (preliminary)

	Sample ID	Total no. of words	A minimum of 3 out of 4 specialists agreed		A minimum of 2 out of 4 specialists agreed	
			Count	%	Count	%
Without new process	EG1	71	6	8%	21	30%
	EG2	69	2	3%	8	12%
	EG3	68	1	1%	6	9%
	EG4	221	2	1%	22	10%
	Overall	429	11	3%	57	13%
With new process	EG5	21	12	57%	16	81%
	EG6	15	13	87%	13	87%
	EG7	35	23	66%	28	80%
	EG8	51	26	51%	38	75%
	Overall	122	74	61%	95	78%

- i. Extend rules for extracting the SAPPhIRE information from technical documents outside of Natural Language Processing (i.e., the parts of speech tags of the words and their syntactical relationship) to include rules based on system modeling techniques (Chakrabarti *et al.*, 2005; Hubka & Eder, 2002). With this change, the category of extracted information is ascertained to determine the most appropriate SAPPhIRE tag.
- ii. An engineering or biological system can have more than one system interaction and currently there is no rule to identify each separately. Hence, the rules need to identify all meaningful system interactions.
- iii. The rules should be able to differentiate a phenomenon from an action and to identify input and state changes if they are not explicit in the natural language description.
- iv. The specialists participating in the intercoder reliability study are familiar with the concepts of the SAPPhIRE modeling. Since the novice users are not familiar with SAPPhIRE modeling, the process should guide them in creating system representation using SAPPhIRE

The first three items described above are related to the reasoning rules for SAPPhIRE information extraction given in Part B of the process. This process was studied previously with the SAPPhIRE specialists. These changes will not impact the process step but will enhance the rules it uses. However, the process or tasks associated with the last bullet item are not defined. This is about representing the extracted information that was not studied before. Therefore, learning and documenting the SAPPhIRE representation process is required first. Hence, work for this second step is carried out in two phases. In the first phase, a simple process called a New Process (Intermediate) is created and a user trial is conducted to observe user cognition. This user trial's observations are then used to define the New Process (Final). The user trial of this first phase also demonstrates whether novice users can follow process instructions for creating representation using SAPPhIRE. Therefore, lessons from this first user trial are used to refine the final user trial at the end.

5.2 User trial procedure

Unlike user trials for the research sub-question 1, users participating in the user trial this time are unfamiliar with the SAPPhIRE

modeling concept. So, before asking them to model, a knowledge session is introduced to train them on understanding systems, systems boundaries, interactions, and definitions of SAPPhIRE constructs. The steps of the User Trial procedure are given in Figure 3. The models created during the user trial are analyzed by comparing them against the reference data set (i.e., information relevant to SAPPhIRE constructs) created by SAPPhIRE specialists. The effectiveness of the process is checked by using the model score. The model score reflects the accuracy and completeness of the SAPPhIRE models created using the process. Sample documents are provided to every participant which are used as use cases during the user trial. These documents are hand-curated to cover information related to every SAPPhIRE construct. The technical information about how the system works, as given in the document, is assumed to be accurate. The intention of the user trial is to check whether a participant can identify the system interactions and build the causal chain for each interaction by identifying the parts, conditions or attributes, external input, governing scientific law associated with that interaction and the state change it produces. Action is the high-level interpretation of a State Change. Hence, it is more important to identify the State Change accurately. The construct Action is kept out of scope as of now in the user trials part of the second step. A training session is conducted to ensure everyone understands the SAPPhIRE concept and the modeling process before creating models.

The choice of a document to be used as a use case becomes critical since the linguistic expressions in the document can be ambiguous and may create interpretational differences due to several factors, such as semantics, writing style, problem context, etc. (Halevy *et al.*, 2009). It is found that designers generate solutions out of multiple analogies by breaking the main design problem into sub-problems. Designers most often search online sources for inspiration to solve the sub-problems (Vattam *et al.*, 2013). Hence, short technical articles with content harvested from online sources are used to validate the process. These short technical articles represent the cases designers typically use during ideation and such short descriptions are about 'how a system works' and, therefore are causal descriptions of systems' working. Technical documents with high school-level science knowledge are used to ensure all user trial participants have the same knowledge of the subject. The background of the participants of the user trials is shown in Sections 5.3 and 5.4.

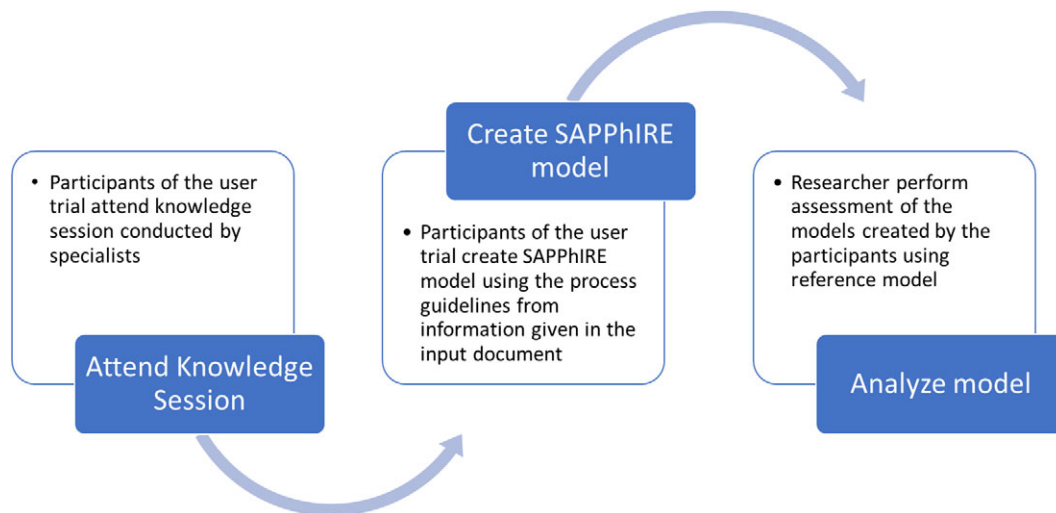


Figure 3. Procedure for the user trial.

5.3 First user trial

In the first phase, a simple process, as shown in Figure 4, is developed to create a SAPPhIRE model from a technical document. As mentioned in Section 5.1, this process is called the New Process (Intermediate). There are no elaborate guidelines for steps 2 & 4 of this process. This process is kept simple with a few steps. Illustrative examples created by specialists are shared with the participants in the user trial. As all the participants have an engineering background, the objective is to watch out how the participants carry out these two steps and identify if there is any common behavioral pattern among the participants. For step 3, the reasoning rules are used and these rules were created earlier (while working on the first step) for extracting SAPPhIRE information and updating, as per the first three items given in Section 5.1. The revised reasoning rules for information extraction are given in Appendix C.

The first user trial was conducted at the Centre for Product Design and Manufacturing department of the Indian Institute of Science. There was a total of 7 students. Participants were from an engineering background without any significant job experience. None of them worked with the SAPPhIRE model before. The participants attended a knowledge session before creating a model. Each participant was given a single use case given in Use Case 1 of Appendix E. This use case is about working of an engineering

system with two main system interactions. This use case was hand-curated and the reference SAPPhIRE model for the use case was created by a SAPPhIRE specialist.

5.4 Results from first user trial

The SAPPhIRE models created by each participant are compared against the reference SAPPhIRE model. The percentage of relevant words identified by each participant for each SAPPhIRE construct is calculated first and then their average. The results show that the average percentage of the relevant words identified by the participants is not high for any of the SAPPhIRE constructs. On closer look, it is found that the participants used systems interaction to decide the number of SAPPhIRE models. We see that there is variation among participants in identifying the number of system interactions (i.e., Phenomena). For example, two participants identified 2 Phenomena, three participants identified 3 Phenomena and two participants identified 4 Phenomena, indicating that interpretation plays a dominating role in deciding the number of SAPPhIRE models. As a result, there is variation in identifying other SAPPhIRE constructs. If a particular Phenomenon is not identified, other associated constructs are also not reported. Results improve significantly when only those two Phenomena are taken, which are

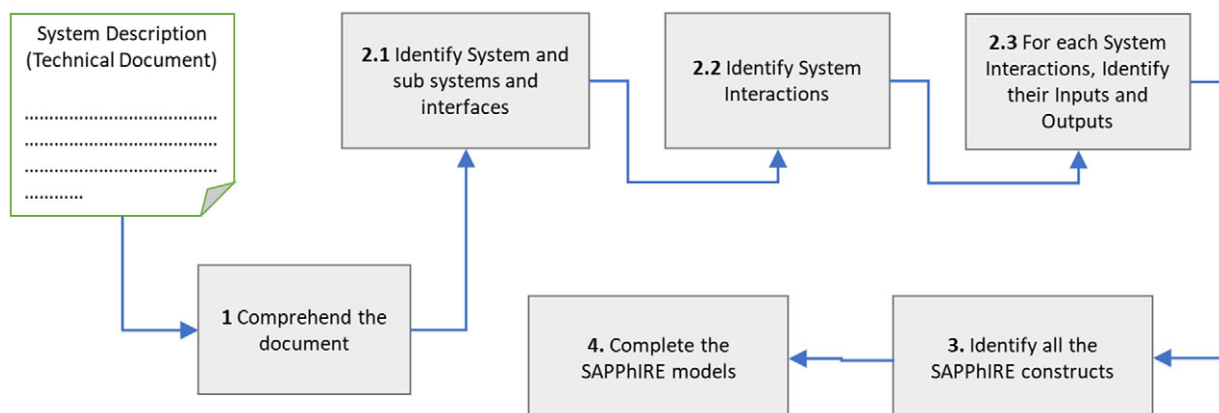


Figure 4. Process diagram of the new process (intermediate) for the first user trial.

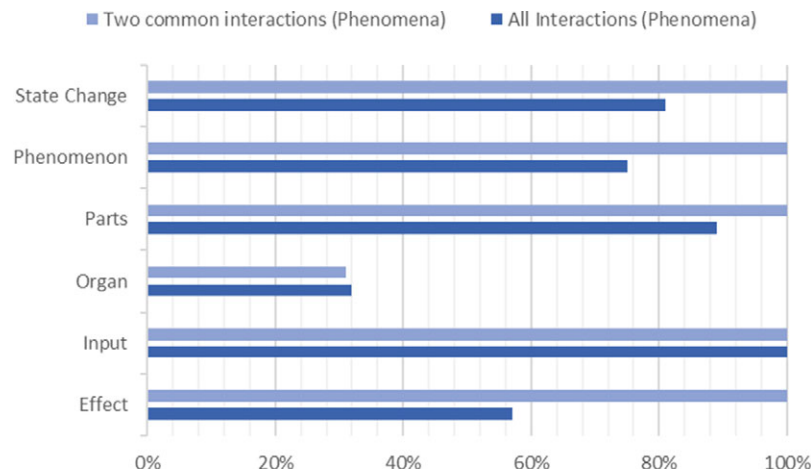


Figure 5. % of the total number of relevant words identified by the participants on average in the first user trial.

common among all the participants, along with the Parts, Organ, Input and State Change associated with these two Phenomena. Figure 5 shows the percentage of the total number of relevant words identified by the participants on average for all four Phenomena and the two common Phenomena.

Based on the analysis of the results and with feedback from the participants, the following lessons from this first user trial are summarized. These are then incorporated into the final user trial procedure:

- System interaction (i.e., Phenomena) can be the basis for the number of SAPPhIRE models. In the absence of a prescriptive guideline, individual interpretation varies in identifying system interactions, resulting in differences in the final model.
- In the first user trial knowledge session, the participants were not clear on what to do (concept) and how to do it (procedure). These two things should be separated.
- The knowledge session should have less presentation and more interactive discussion. Due to limited interactions during the knowledge session in the first user trial, participants did not get opportunities to clarify doubts about the SAPPhIRE concept and process guidelines.
- The first user trial process was heavy on the use of work product templates and the purpose of these templates was unclear to the participants.
- Except for the Organ, in general, all the words identified by the participants are relevant for the SAPPhIRE constructs, and very few words were not relevant for the modeling. This implies that when the process is applied correctly, it produces relevant results.

5.5 Final user trial

Based on the lessons learned from the first user trial, a prescriptive process was created for information extraction and representation using the SAPPhIRE model. As mentioned in Section 4.1, this process is called the New Process (Final). The process diagram of the revised process is shown in Figure 6. An illustration of the process is given in Appendix F using the example of the working of a Solenoid Valve, whose description is also given in the same Appendix.

The number of work product templates is reduced to one template in the final user trial. For the knowledge session, interactive discussion is emphasized over presentation. The SAPPhIRE concept

is first explained, and then, the modeling process is taught before asking the participants to generate the SAPPhIRE representation.

The SAPPhIRE Modeling table to organize the extracted information is given in Appendix D.

Summary of Rules at a high level

- Phenomena is a technical process inside the system indicated by Action Verbs. It involves the transformation or transfer of material or energy or information, which are the subject and object of those verbs
- Parts are the Material Nouns representing material entities of the system
- Organs are the properties and conditions required for the interaction to happen, represented by prepositions, conjunctions, or adverbs
- Input is the material, energy, or information that comes from outside the boundary, represented as the subjects and objects of the action verbs
- State change is the change of Property of the entity and its surroundings involved in the interaction, represented as the subjects and objects of the action verbs

The user trial was conducted at the Centre for Product Design and Manufacturing department of the Indian Institute of Science. There was a total of 16 student participants. Participants come from an Engineering background. All the participants were fresh from college without any job experience. None of them worked with the SAPPhIRE model before. Participants attended a training session before creating the model. A question-and-answer session was allowed before modeling to clarify any doubts. At the end of the user trial, participants' feedback was collected for (a) the level of clarity of the SAPPhIRE concept before modeling (b) the level of clarity on the modeling process before modeling and (c) the amount of time allocated for the modeling.

Each participant was given the same set of two use cases for modeling given in Appendix E, Use Case 1, and Use Case 2. Each use case is about the working of an engineering system with two main system interactions. These use cases are hand-curated. The reference data set (i.e., information relevant to SAPPhIRE constructs) for each use case is created by a SAPPhIRE specialist and validated by two other SAPPhIRE specialists. The use cases and their reference SAPPhIRE data set, created and validated by the specialists, are given in Appendix E.

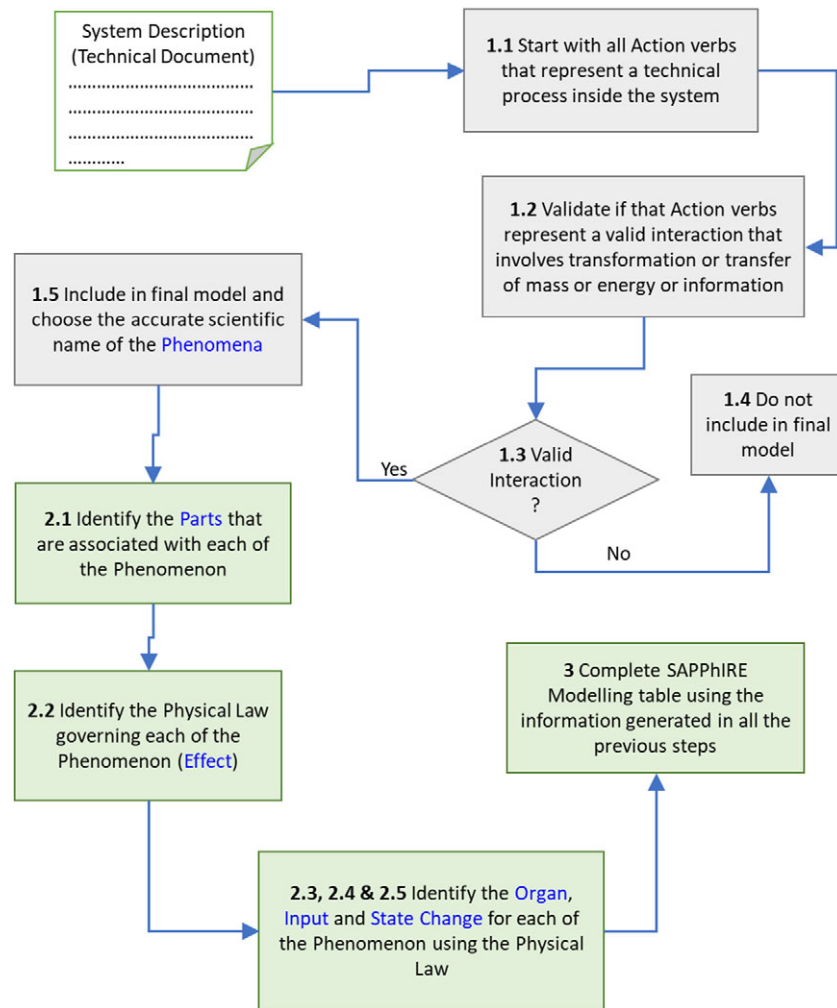


Figure 6. Process diagram of the new process (final) for the final user trial.

5.6 Results of the final user trial

The following metrics are used to analyze the models created by the participants.

Approach A (Level of Agreement with the reference data set): This approach is based on the number of participants who correctly identified a SAPPhIRE construct. For each relevant word/clause in the reference data set, calculate the agreement score in terms of percentage as:

$$\% \text{Agreement} = \frac{\text{No. of participants who identified that word as relevant}}{\text{Total No. of participants}} \times 100 \quad (\text{i})$$

Then, calculate the agreement of each SAPPhIRE construct based on all the relevant words/clauses belonging to that SAPPhIRE construct. This % Agreement score indicates, on average, how many participants identified the words/clauses determined as relevant by a specialist. The % Agreement score, therefore, implies repeatability since higher agreement indicates many participants can apply the rules consistently.

Approach B (Accuracy): This approach is based on Precision and Recall scores (Russel & Norvig, 2010). Here, first, calculate each participant's Precision and Recall for each SAPPhIRE

construct. Then, calculate the average Precision and Recall of all the participants. A precision value indicates how many of the identified Words/Clauses are relevant. A recall value indicates how many of the relevant words/clauses are identified by the participant.

$$\text{Precision} = \frac{\text{No. of True Positives}}{\text{No. of True Positives} + \text{No. of False Positives}} \quad (\text{ii})$$

$$\text{Recall} = \frac{\text{No. of True Positives}}{\text{No. of True Positives} + \text{No. of False Negatives}} \quad (\text{iii})$$

$$\text{F1 score} = \frac{2 * \text{Precision} * \text{Recall}}{(\text{Precision} + \text{Recall})} \quad (\text{iv})$$

The terms True Positives, False Positives, and False Negatives used in the equations (ii) (iii), and (iv) are defined in Table 3.

Figure 7 shows the results of the analysis of models using Approach A (equation i) and Approach B (equation ii-iv). Figure 8 shows the number of system interactions identified by the participants. Figure 9 shows % Agreement scores for use case 1 and use case 2 separately, mainly to check for any noticeable differences between them.

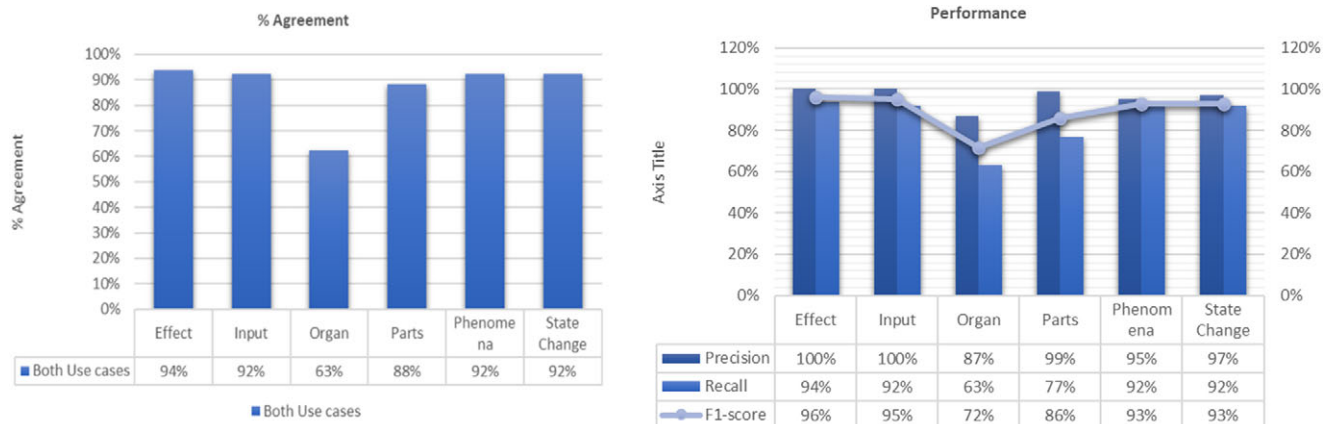
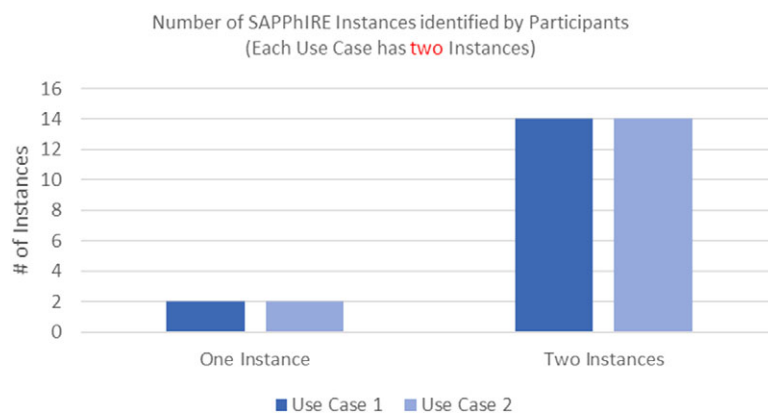
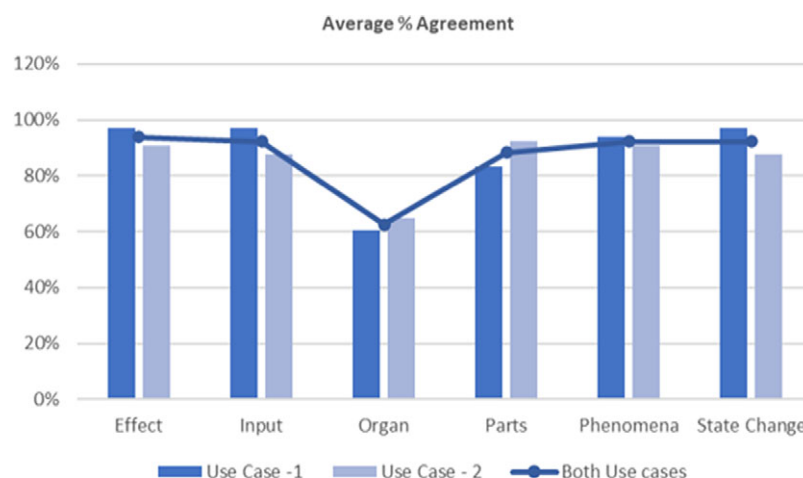
Table 3. Definitions of terms true positives, false positives and false negatives

Actual (reported by specialists)	By the user trial participants	
	Reported in result set	Not reported in result set
Relevant	True positive	False negative
Not relevant	False positive	True negative

5.7 Discussions of results of the final user trial

The following observations are made from the final user trial:

- The average Agreement score of 87% and the Average F1-Score of 89% are in the same order. 87.5% of participants reported two instances correctly for both use cases.
- In general, it is found that precision is higher than the recall values. High precision means lesser numbers of irrelevant

**Figure 7.** Analysis of models in the final user trial.**Figure 8.** Number of system interactions identified by the participants in the final user trial.**Figure 9.** % Agreement scores for use case 1 and use case 2 in the final user trial.

words/clauses (i.e., More True Positives and fewer False Positives). Out of 16 participants, 13 of them on average, could identify the 32 Relevant words/clauses (determined by specialists). On the other hand, there were 10 words/clauses which are not Relevant but only 1.5 participants on average marked them as Relevant. This implies that the participants can follow the process in most situations.

- However, lower recall values indicate that a few participants did miss identifying relevant information. This could be due to the lack of clarity in the input natural language documents and the knowledge gap of some participants.
- Except for Organ, precision values are higher than 95% for the other SAPPhIRE constructs. However, the recall values are lesser. This is primarily because some participants missed identifying one or two Phenomena and associated other SAPPhIRE constructs like the Effects, Parts, Inputs, and State Changes of the linked Phenomena (e.g., two participants missed to identify Blackbody radiation as the Phenomenon in the first use case and two participants missed to identify Volume Expansion due to heating as the Phenomenon in the second use case).
- In both use cases, a few parts are mentioned whose causal role is not given in the document. Since these material entities represent a system object, they were already included in the reference data set. Some participants reported them as Parts, while others ignored them. Hence, a reduction in recall values for the Parts is observed.
- Precision and recall values for the Organ are much lower than the other SAPPhIRE constructs. This could be due to a knowledge gap. It is noticed that most participants overlooked the condition for system interaction given in the document. E.g., in the case of Heat Transfer, the temperature gradient is a condition for the Phenomenon to occur. But except for one participant, the rest overlooked it. Similarly, a lightbulb emits light only when the filament temperature is high enough. Though this was mentioned in the document, a majority of the participants overlooked it. 50% of the false positive information is related to Organ. These would be investigated by improving the rules and guidelines around Organ in the next revision.
- The observations from the final user trial are used to validate the rules, which would be useful for developing automation support later. It is found that the rules work without any changes in 93.75% of the 32 relevant (true positive) words/clauses. However, rules need to be updated in 6.25% of the 32 relevant (true positive) words/clauses. In these cases, rules did not work but participants used their domain knowledge in heat transfer to identify the relevant words for the SAPPhIRE construct. Also, 10 words/clauses are not relevant but participants falsely interpret the word/clause as a true SAPPhIRE construct (False Positive). These cases would be looked at in the future.
- Researchers previously reported function-based reasoning as a pervasive and important basis for conceptual designs (Umeda and Tomiyama, 1997; Mulet et al., 2009; Escrig et al. 2009). In this research, it is observed that identifying system interactions at the beginning and then selecting each system interaction as a Phenomenon to use as the basis for finding other SAPPhIRE constructs aligns with the cognition of designers and leads to better consistency and accuracy of information extraction. If 'Function' is considered as what the

system is intended to do or produce (Chakrabarti, 1998), it then points to the interaction or interactions a system has with its surroundings. This might be a major reason as to why the process led to better consistency.

6. Conclusions and future work

'Multiplicity', which is about maintaining diversity and variety of content, is an important requirement of any Design-by-Analogy database (Srinivasan et al., 2010; Yargin et al., 2015). Therefore, a large database to enable Design-by-Analogy with systems models, like the SAPPhIRE model, will help have more (count) inspirations or stimuli in the conceptual phase of design, and therefore can potentially support creativity. Design support like IDEA-INSPIRE and DANE accurately represents systems knowledge using its multi-level ontological schema, called the SAPPhIRE and the SBF models. However, such databases have a limited number of models and creating new models is effort intensive (Willcox and Dufloy, 2023). But natural language documents are available in plenty, for e.g., Espacenet (<https://worldwide.espacenet.com/>) alone has free access to over 140 million patent documents from multiple countries. Many such Patent databases (e.g., USPTO, PatentScope, DEPATISnet, etc.) have vast numbers of documents. Hence, it is believed that these large number of documents can serve as a very good source of inspiration if the knowledge of these documents is converted into a structured format using the SAPPhIRE model (in IDEA-INSPIRE) and made available to the designers for inspiration. Currently, this document conversion process into a model is not well understood. It is very effort and time-intensive, and the outcome of this process heavily depends on the expertise of the person trying to develop the ontological model from a technical document. Not much research has been reported on understanding the process and its end-to-end performance.

This research adopted a user-centric, iterative approach for creating a simplified and generalized process with a set of rules for creating an accurate and repeatable ontology-based data model from natural language descriptions. This process is aligned with the function-based design reasoning and extracts information for SAPPhIRE constructs, ensuring the context of the document instead of individual sentences. The process starts with identifying system interactions given in the document by taking clues from the action verbs given in the document. Then, it builds the causal chain for each system interaction by extracting relevant information from the subject and predicate of the sentences with Action verbs. The consistency (i.e., many users produce the same results) and accuracy (i.e., every user produces correct results) of the final output of the process are reported in this research, indicating the end-to-end process performance.

Work started with the first step on understanding the current process by conducting an intercoder reliability study with the SAPPhIRE specialists. This intercoder reliability study, without any intervention, indicated differences among the specialists about the SAPPhIRE information extracted from a document described in natural language. It is only for 3% of the words/clauses where 3 out of 4 users agreed with the extracted information. Results were analyzed and it was found that this was due to the interpretational differences and not applying the SAPPhIRE definitions consistently to the context. Based on input from the workshop with specialists after the first intercoder reliability study, the New Process (Preliminary) was developed and another intercoder reliability

study was conducted with the same specialists, where the score improved to 61% from 3%.

The second step is to find out what a new process should be. The SAPPhIRE specialists' cognition is captured in the New Process (Preliminary) and then generalized and simplified for all users through multiple iterations to come up with the New Process (Final). The results of the final user trial with this New Process (Final) show significant performance with the % Agreement score (consistency) at 87% and the F1-score score at 89% (accuracy). The previously reported F1-score for classifying target words or clauses to a SAPPhIRE label was 70% (Keshwani & Chakrabarti, 2017) and to an SBF label was 84% (Goel et al., 2020). These numbers represent the classification task's accuracy and do not include data preparation tasks before classification. The accuracy reported in this research is the accuracy of the end-to-end process. The observations regarding Organ need to be looked into to improve its scores in the future.

Technical documents in natural languages can be written in various styles with different levels of technical details. Future research will cover these areas (not an exhaustive list): processing multiple descriptions of the same system, large-size sample cases for checking meaningfulness or logical sense (syntactical analysis), and developing a support tool for automating the process.

Supplementary material. The supplementary material for this article can be found at <http://doi.org/10.1017/S0890060424000118>.

References

- Baldussu A, Cascini G, Rosa F and Rovida E (2012) Causal models for bio-inspired design: A comparison. In *DS 70: Proceedings of DESIGN 2012, the 12th International Design Conference, Dubrovnik, Croatia*, pp. 717–726.
- Bhattacharya K, Bhatt AN, Ranjan BSC, Keshwani S, Srinivasan V and Chakrabarti A (2023) Extracting information for creating SAPPhIRE model of causality from natural language descriptions. *Design Computing and Cognition* 22, 3–20. https://doi.org/10.1007/978-3-031-20418-0_1
- Blessing LT and Chakrabarti A (2009) *DRM: A Design Research Methodology*. Springer London, pp. 13–42.
- Bonnardel N (2000) Towards understanding and supporting creativity in design: Analogies in a constrained cognitive environment. *Knowledge-Based Systems* 13(7–8), 505–513.
- Cascini G and Zini M (2011) Computer-aided comparison of thesauri extracted from complementary patent classes as a means to identify relevant field parameters. In *Global Product Development*. Berlin, Heidelberg: Springer, pp. 555–566. https://doi.org/10.1007/978-3-642-15973-2_56
- Chakrabarti A (1998) Supporting two views of function in mechanical designs. In *Proceedings 15th National Conference on Artificial Intelligence, AAAI (Vol. 98)*, pp. 26–30.
- Chakrabarti A (2009) The art and science behind successful product launches. In Chapter 2 Raghavan NRS and Cafeo JA (eds.), *Product Research*. <https://doi.org/10.1007/978-90-481-2860-02>
- Chakrabarti A, Sarkar P, Leelavathamma B and Nataraju BS (2005) A functional representation for aiding biomimetic and artificial inspiration of new ideas. *Artificial Intelligence for Engineering Design, Analysis and Manufacturing* 19(2), 113–132.
- Chen X, Jia S and Xiang Y (2020) A review: Knowledge reasoning over knowledge graph. *Expert Systems with Applications* 141, 112948.
- Cheong H and Shu LH (2012) Automatic extraction of causally related functions from natural-language text for biomimetic design. In *International Design Engineering Technical Conferences and Computers and Information in Engineering Conference (Vol. 45066)*, pp. 373–382). American Society of Mechanical Engineers.
- Chulvi V, Mulet E, Chakrabarti A, López-Mesa B and González-Cruz C (2012) Comparison of the degree of creativity in the design outcomes using different design methods. *Journal of Engineering Design* 23(4), 241–269.
- Deldin JM and Schuknecht M (2014) The AskNature database: Enabling solutions in biomimetic design. In *Biologically Inspired Design*. London: Springer, pp. 17–27.
- Escrig EM, Chakrabarti A, Chulvi V, Mesa BL, González MC, Porcar A and Galán J (2009) How strongly does functional analysis influence creativity? In *Comunicaciones presentadas al XIII Congreso Internacional de Ingeniería de Proyectos: celebrado en Badajoz el 8, 9 y 10 de julio 2009*. Asociación española de ingeniería de proyectos (AEIPRO), p. 49.
- Fu K, Moreno D, Yang M and Wood KL (2014) Bio-inspired design: An overview investigating open questions from the broader field of design-by-analogy. *Journal of Mechanical Design* 136(11).
- Goel A, Hagopian K, Zhang S and Rugaber S (2020) Towards a virtual librarian for biologically inspired design. In *Proceedings of the 9th International Conference on Design Computing and Cognition*, pp. 377–396.
- Goel A, Rugaber S and Vattam S (2009) Structure, behavior & function of complex systems: The SBF modeling language. *Artificial Intelligence in Engineering Design, Analysis and Manufacturing* 23(1), 23–35. <https://doi.org/10.1017/S0890060409000080>
- Goel AK (1997) Design, analogy, and creativity. *IEEE Expert* 12(3), 62–70.
- Goel AK and Shu LH (2015) Analogical thinking: An introduction in the context of design. *Artificial Intelligence in Engineering Design, Analysis and Manufacturing* 29(2), 133–134. <https://doi.org/10.1017/S0890060415000013>
- Gonçalves M, Cardoso C and Badke-Schaub P (2016) Inspiration choices that matter: The selection of external stimuli during ideation. *Design Science* 2. <https://doi.org/10.1017/dsj.2016.10>
- Halevy A, Norvig P and Pereira F (2009) The unreasonable effectiveness of data. *IEEE Intelligent Systems* 24(2), 8–12. <https://doi.org/10.1109/mis.2009.36>
- Han J, Sarica S, Shi F and Luo J (2022) Semantic networks for engineering design: State of the art and future directions. *Journal of Mechanical Design* 144(2).
- Han J, Shi F, Chen L and Childs P (2018) A computational tool for creative idea generation based on analogical reasoning and ontology. *AI EDAM* 32(4), 462–477. doi:<https://doi.org/10.1017/S0890060418000082>
- Hashemi Farzaneh H, Helms MK and Lindemann U (2015) Visual representations as a bridge for engineers and biologists in bio-inspired design collaborations. In *International Conference on engineering Design, ICED15*.
- Helms M, Vattam SS and Goel AK (2009) Biologically inspired design: Process and products. *Design Studies* 30(5), 606–622.
- Howard TJ, Dekoninck EA and Culley SJ (2010) The use of creative stimuli at early stages of industrial product innovation. *Research in Engineering Design* 21(4), 263–274. <https://doi.org/10.1007/s00163-010-0091-4>
- Hsieh SH, Lin HT, Chi NW, Chou KW and Lin KY (2011) Enabling the development of base domain ontology through extraction of knowledge from engineering domain handbooks. *Advanced Engineering Informatics* 25(2), 288–296. <https://doi.org/10.1016/j.aei.2010.08.004>
- Hubka V and Eder WE (2002) Theory of technical systems and engineering design synthesis. In *Engineering Design Synthesis*. London: Springer, pp. 49–66. https://doi.org/10.1007/978-1-4471-3717-7_4
- Jagtap S (2018) Design creativity: Refined method for novelty assessment. *International Journal of Design Creativity and Innovation* 7(1–2), 99–115. <https://doi.org/10.1080/21650349.2018.1463176>
- Jiang S, Hu J, Wood KL and Luo J. (2022) Data-Driven Design-By-Analogy: State-of-the-Art and Future Directions. *Journal of Mechanical Design* 144(2)
- Keshwani S (2018) Supporting designers in generating novel ideas at the conceptual stage using analogies from the biological domain, PhD Thesis, IISc
- Keshwani S and Chakrabarti A (2015) Influence of analogical domains and abstraction levels on novelty of designs. In *DS79: Proceedings of the Third International Conference on Design Creativity*. Indian Institute of Science, Bangalore.
- Keshwani S and Chakrabarti A (2017) Detection and splitting of constructs of SAPPhIRE model to support automatic structuring of analogies. In *DS 87–4 Proceedings of the 21st International Conference on Engineering Design (ICED 17) Vol 4: Design Methods and Tools*. Vancouver, Canada, pp. 603–612.
- Kim KY and Kim YS (2011) Causal design knowledge: Alternative representation method for product development knowledge management. *Computer-Aided Design* 43(9), 1137–1153.

- Lee HS and Holyoak KJ (2008) The role of causal models in analogical inference. *Journal of Experimental Psychology: Learning, Memory, and Cognition* 34(5), 1111.
- Lenau TA, Keshwani S, Chakrabarti A and Ahmed-Kristensen S (2015) Bio cards and level of abstraction. In *International Conference On Engineering Design, ICED15*
- Li Z and Ramani K (2007) Ontology-based design information extraction and retrieval. *Artificial Intelligence Engineering Design Analysis and Manufacturing* 21(2), 137–154. <https://doi.org/10.1017/s0890060407070199>
- Moreno DP, Blessing LT, Yang MC, Hernández AA and Wood KL (2016) Overcoming design fixation: Design by analogy studies and nonintuitive findings. *AI EDAM* 30(2), 185–199.
- Mulet E, Chakrabarti A, Chulvi V, López-Mesa B, González M, Porcar A and Galán J (2009) How Strongly Does Functional Analysis Influence Creativity?. *Journal of Mechanical Design* 136(10).
- Murphy J, Fu K, Otto K, Yang M, Jensen D and Wood K (2014) Function based design-by-analogy: A functional vector approach to analogical search. *Journal of Mechanical Design* 136(10).
- Nagai Y and Taura T (2015) Studies of design creativity: A review and its prospects. *Journal of the Indian Institute of Science* 95(4), 341–351.
- Nagel JK, Nagel RL, Stone RB and McAdams DA (2010) Function-based, biologically inspired concept generation. *Artificial Intelligence for Engineering Design, Analysis and Manufacturing: AI EDAM* 24(4), 521.
- Nédellec C, Nazarenko A and Bossy R (2009) Information extraction. In *Handbook on Ontologies*. Berlin, Heidelberg: Springer, pp. 663–685. https://doi.org/10.1007/978-3-540-92673-3_30
- Russell JS and Norvig P (2010) *Artificial Intelligence A Modern Approach*, 3rd Edn. Prentice Hall.
- Salvatore TM (2003) Data modeling: Entity-relationship data model. In Bidgoli H (ed.), *Encyclopaedia of Information Systems*. Elsevier, pp. 489–503. <https://doi.org/10.1016/B0-12-227240-4/00034-4>.
- Sarica S, Han J and Luo J (2023) Design representation as semantic networks. *Computers in Industry* 144, 103791.
- Sarica S and Luo J (2021) Design knowledge representation with technology semantic network. *Proceedings of the Design Society* 1, 1043–1052.
- Sarkar P and Chakrabarti A (2008) The effect of representation of triggers on design outcomes. *Artificial Intelligence for Engineering Design, Analysis and Manufacturing: AI EDAM* 22(2), 101.
- Sartori J, Pal U and Chakrabarti A (2010) A methodology for supporting transfer in biomimetic design. *AI EDAM* 24(4), 483–506.
- Shu LH (2010) A natural-language approach to biomimetic design, *Artificial Intelligence for Engineering Design, Analysis and Manufacturing* 24, 507–519
- Siddharth L, Blessing L, Wood KL and Luo J (2022) Engineering Knowledge Graph from Patent Database. *Journal of Computing and Information Science in Engineering* 22(2).
- Siddharth L, Chakrabarti A and Venkataraman S (2018) Representing complex analogues using a function model to support conceptual design. In *International Design Engineering Technical Conferences and Computers and Information in Engineering Conference* (Vol. 51739, p. V01BT02A039). American Society of Mechanical Engineers.
- Simon HA (1969) *The Sciences of the Artificial*. Cambridge: MIT Press.
- Song H and Fu K (2019) Design-by-analogy: Exploring for analogical inspiration with behavior, material, and component-based structural representation of patent databases. *Journal of Computing and Information Science in Engineering* 19(2).
- Srinivasan V and Chakrabarti A (2009) SAPPhIRE—an approach to analysis and synthesis. In *DS 58–2: Proceedings of ICED 09, the 17th International Conference on Engineering Design, Vol. 2, Design Theory and Research Methodology*, Palo Alto, CA, USA, 24–27.08.2009, pp. 417–428)
- Srinivasan V and Chakrabarti A (2010) Investigating novelty–outcome relationships in engineering design. *AI EDAM* 24(2), 161–178.
- Srinivasan V, Chakrabarti A and Lindemann U (2012) A framework for describing functions in design. In *DS 70: Proceedings of DESIGN 2012, the 12th International Design Conference*, Dubrovnik, Croatia.
- Srinivasan V, Chakrabarti A and Lindemann U (2015) An empirical understanding of use of internal analogies in conceptual design, *Artificial Intelligence for Engineering Design, Analysis and Manufacturing (AI EDAM)* 29(2), 147–160. <http://doi.org/10.1017/S0890060415000037>
- Srinivasan V, Song B, Luo J, Subburaj K, Elara MR, Blessing L and Wood K (2017) Investigating effects of stimuli on ideation outcomes. In *DS 87–8 Proceedings of the 21st International Conference on Engineering Design (ICED 17) Vol 8: Human Behaviour in Design*, Vancouver, Canada, 21–25.08.2017, pp. 309–318.
- Umeda Y and Tomiyama T (1997) Functional reasoning in design. *IEEE Expert* 12(2), 42–48.
- Vandevenne D, Pieters T and Dufloy JR (2016) Enhancing novelty with knowledge-based support for Biologically Inspired Design. *Design Studies* 46, 152–173.
- Vandevenne D, Verhaegen PA, Dewulf S and Dufloy J (2016b) SEABIRD: scalable search for systematic biologically inspired design. *Artificial Intelligence for Engineering Design, Analysis and Manufacturing* 30(1), 78–95.
- Vattam S and Goel A (2013) Seeking bioinspiration online: A descriptive account. In *DS 75–7: Proceedings of the 19th International Conference on Engineering Design (ICED13), Design for Harmonies, Vol. 7: Human Behaviour in Design*, Seoul, Korea.
- Vattam S, Wiltgen B, Helms M, Goel AK and Yen J (2011) DANE: Fostering creativity in and through biologically inspired design. In *Design Creativity 2010*. London: Springer, pp. 115–122.
- Verhaegen PA, D’hondt J, Vandevenne D, Dewulf S and Dufloy JR (2011) Identifying candidates for design-by-analogy. *Computers in Industry* 62(4), 446–459.
- Waldmann, Michael R., and York Hagmayer (2013) Causal reasoning. In Reisberg D (ed.), *The Oxford Handbook of Cognitive Psychology*, Oxford Library of Psychology. Online edn, Oxford Academic, 3 June, <https://doi.org/10.1093/oxfordhb/9780195376746.013.0046>, accessed 29 October 2023.
- Wilcox M, Ayali A and Dufloy JR (2020) Where and how to find bio-inspiration? A comparison of search approaches for bio-inspired design. *CIRP Journal of Manufacturing Science and Technology* 31, 61–67.
- Wilcox M and Dufloy JR (2023) Free-text inspiration search for systematic bio-inspiration support of engineering design. *AI EDAM* 37, e21.
- Wimalasuriya DC and Dou D (2010) Ontology-based information extraction: An introduction and a survey of current approaches. *Journal of Information Science* 36(3), 306–323.
- Yang J, Kim E, Hur M, Cho S, Han M and Seo I (2018) Knowledge extraction and visualization of digital design process. *Expert Systems with Applications* 92, 206–215. <https://doi.org/10.1016/j.eswa.2017.09.002>
- Yargin GT and Crilly N (2015) Information and interaction requirements for software tools supporting analogical design. *AI EDAM* 29(2), 203–214.
- Zhang C, Zhou G, Chang F and Yang X (2020) Learning domain ontologies from engineering documents for manufacturing knowledge reuse by a biologically inspired approach. *The International Journal of Advanced Manufacturing Technology* 106(5), 2535–2551. <https://doi.org/10.1007/s00170-019-04772-1>