

Mental Models Report

By Len Kromkamp, Marieke Voorhuijzen, Tom van 't Westeinde

Contents

Process

- Mental Models behind Intelligent Systems	3
- Correlation User Feedback, System Intelligence and Mental Model	4
- Intelligent Systems on Individualism in Modern Society	5
- Implications of Intelligent Systems in Reality	6
- Understanding Fatality et al.	7

Conclusions

- Mutual Intelligibility	8
- Pre-Knowledge and the Mental Model Cycle	9
- Expectation and Capability on Satisfaction	10
- Intelligence versus Mental Model Complexity	11

Process

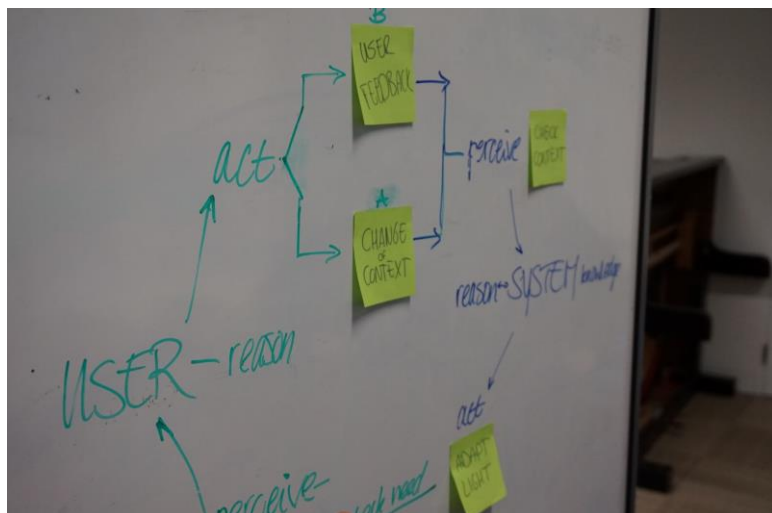
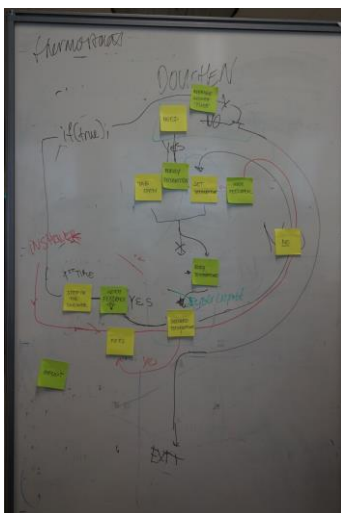
Mental Models behind Intelligent Systems

At the start of the module we were introduced into the basic terminology of mental models. During this lecture we were made familiar with what a mental model is by getting quick examples, but also doing small exercises (call someone on your smartphone).

This initial, albeit limited, information gave us the opportunity to decide on five examples of intelligent systems within five different domains:

- Home: The Smart Shower, automatic temperature and shower duration settings based on the user.
- Work: Contextual Adaptive Lighting System within the office space to create supportive lighting based on the context of the user (think: meeting, brainstorm session, break, co-creation, exploration, research, report writing, etcetera).
- Mobility: The New OV-Chip Check-In, a new way of checking in with your public transport pass, along with others to automatically travel together in order to get the price reduction from different subscriptions (40% reduction, Dal-uren, etcetera).
- Automotive: Personalized Driving Behavior within autonomous cars, in which the user can correct the system in order to learn it different driving patterns that represent the driving style of the user.
- Communication: Contextual Notification Hierarchy in which the smart phone, or any social smart device, decides when and how to let through notifications based on context (current surroundings of the user) as well as the sender of the notification (relative, friend, or maybe just someone random) and the application which sends the notification (phone, text message, Facebook, Angry Birds).

We developed several flowcharts of mental models based on these small concepts. By encompassing the information we got during the lecture, implementing this in the mental models and actively discussing what we were creating allowed us to start getting a better understanding of what a mental model is, how it is relevant to a designer but also how impressively immense it can be.

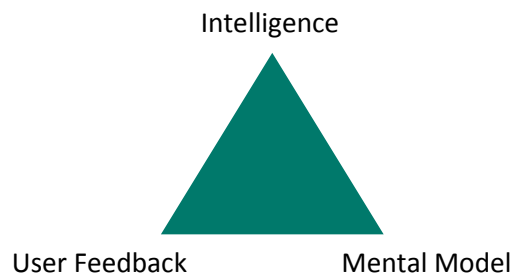


Process

Correlation User Feedback, System Intelligence and the Mental Model

The mental models we generated, were eventually all based on the principle of 'Perceive-Reason-Act' as explained in the lectures. This principle works on both sides of an interaction: user and system. The act of the user is perceived by the system, which reasons his logical outcome and acts, which then again is perceived by the user, who will then reason to act.

From this we generated a triangle-model in which we placed the main influences of an interaction:

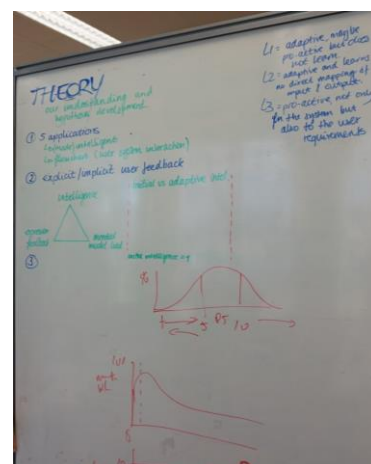
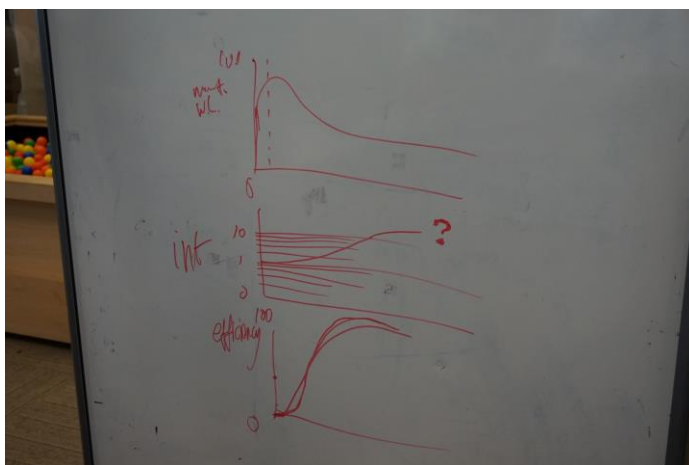


The specific meaning of each terms are as follows:

- Intelligence, of the system.
- User Feedback, explicit and within a specific context towards the system.
- Mental Model, the complexity of the user's mental model on the system.

How this works in an interaction is as follows: when the intelligence of an example system would be high enough to enable to user to be free in its interaction with the system, there is no explicit feedback from the user needed to let the interaction work, the mental model complexity of the user will be lower than a system which incorporates lower intelligence which does require explicit feedback in order to function properly.

This is most definitely a proposition we felt comfortable making at this point in the module. It is important to point out that several aspects of mental models are not yet implemented in this hypothesis.



Process

Intelligent Systems on Individualism in Modern Society

During the previous century we've seen an increase in individualism in societies. This really can be seen when looking at society from the perspective of the Maslow hierarchy of needs.

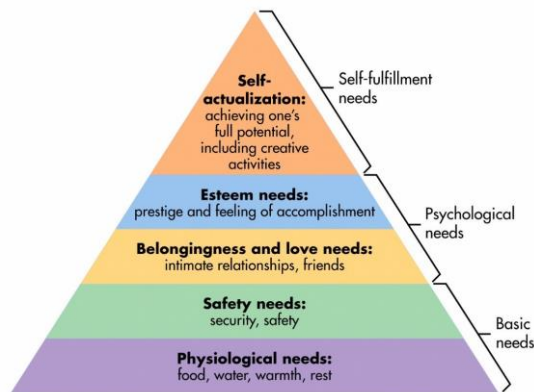
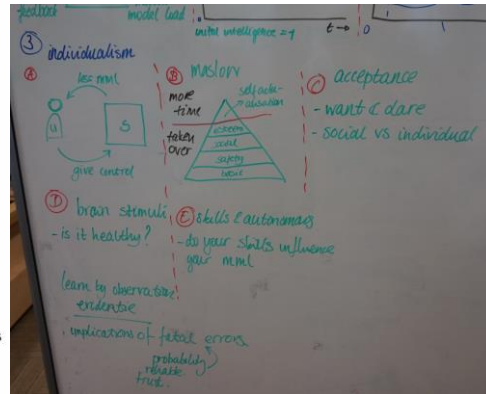


Image credit: 21st Century Tech



Especially in Western countries (West-Europe, The United States of America) societies have been built that have developed the bottom needs in an extensive and complement way. In a culture where basic needs and psychological needs have developed into a standard, the opportunity for self-fulfillment becomes more needed and accessible.

This is partly due to the fact that other intelligent systems take over basic tasks. A good example in our opinion is the supermarket. It may not be fully digital or automated, but in a sense it actually is. It is a big system in which all parties carry out their tasks from getting a product from their origin into the supermarket. The consumer can then buy the product, without ever having to come into contact with the process. Image that in today's society someone would need to go out and hunt for a deer so they can have dinner that night. It is a weird idea, but not long ago this was reality which now has been automated.

The user gives out control to an external party to fulfill a specific need of the user (in the example the basic need for food) and while doing so the user becomes inattentive to the process. You could say the mental model of the user towards the usage of the product has become less complex, becomes so many factors have been externalized.

An important aspect in the course of the implementation of such organization is the societal and individual acceptance towards said organization. Does the (end) user want to give away the control or do they even care?

Process

The Implications of Intelligent Systems in Reality

When analyzing our current hypotheses and with the expert view of Jacques we concluded that, in order for us to get a more in-depth understanding of the implications of mental models in intelligent systems, we had to get more knowledge about existing mental models and their opportunities and problems instead of only theorize about such models.

How does the system act in specific situation? How does the user's knowledge influence the interaction between the user and the system? How does it influence the feedback of both parties and how does this influence the trust the user has with the system?

A great part of the acquainted knowledge came from the "Learning from a Learning Thermostat" paper by Rayoung Yang and Mark Newman (see appendices). The research from the paper is based on an extensive user experiment in which participants use the Nest smart thermostat.

From the paper it was apparent that the smart thermostat is actually pretty flawed. Many users reported (perceived) system errors, but also developed workarounds when encountering such errors. While these errors can be partially attributed to the system, the main problem was communication between system and user, and vice versa.

Specifically in the case of the Nest, the system wasn't always able to understand the intention behind senses behavior, while the user had difficulty in understanding how the system works and what it is doing.

Notable examples are:

- The system didn't understand if user feedback was routine behavior or temporary behavior.
- The user makes assumptions about the expected system behavior, these assumptions are not or only partially met.
- Participants are disappointed by lack of intelligent behavior.
- System does not communicate its behavior in an explicit way, which made it hard for participants to understand it.

At this point we concluded that not necessarily all the actions need to be communicated from and with either party, but more importantly that they need to understand each other where necessary. The intelligibility of the system towards the user needs to be apparent, not continual of necessity but discernible when needed, while at the same time the user should be able to see if their own actions and input are perceived befittingly by the system.



Process

Understand Fatality et al.

An important part of a system is if it fulfills the need of the user consistently and successfully. How impactful is the occurrence of fatality in a system and how does it influence the mental model?

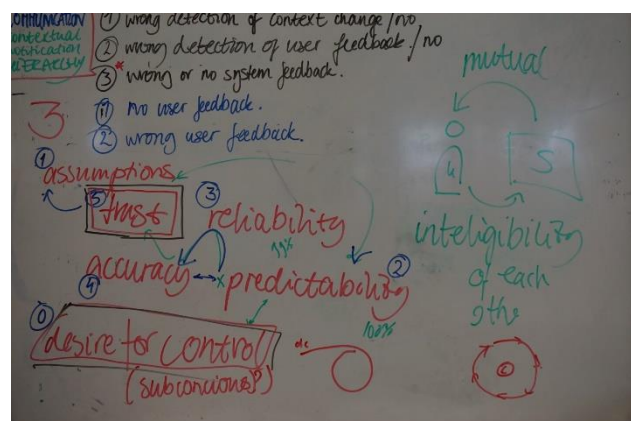
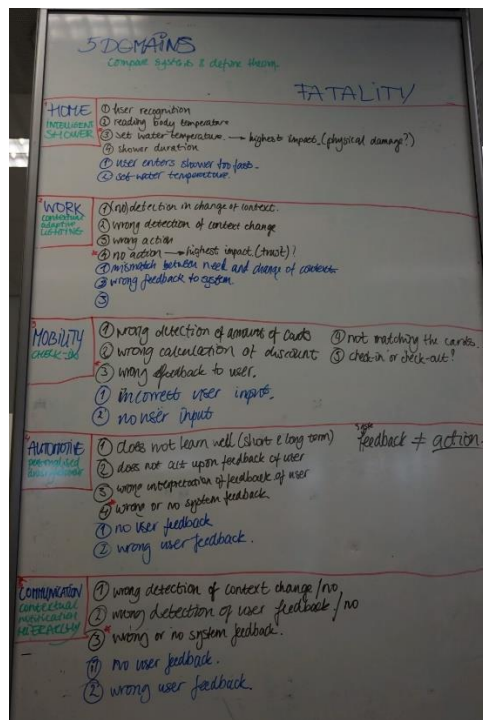
To understand fatality better, we described as much failures of our domain examples and pointed out to what extend the biggest failure impacted the user. By biggest failure we already chose the failure that we perceived as most impactful in the domain.

It became apparent that the majority of the failures aren't necessary horrible, but more likely perceived as annoying. An example of this is the Contextual Adaptive Lighting System, in which the biggest failure would be the system not changing the light when there's a context change, either by being unable to detect the context change or simply refusing to change the lights. In this case the impact of the failure could even be unwitnessed and in a worst case scenario the users would notice and would be bothered by it, perhaps needing to change it manually and learn the system that it made a mistake.

At the same time a fatality error could almost be fatal, or in the least harmful to physical state of the user. An example of this is the Smart Shower. If at any point the system perceives information wrongly or makes a wrong decision and sets the temperature too high (assuming there's no temperature lock), it could scald the user's skin.

As previously stated, it is important for both sides of the interaction to understand each other: mutual intelligibility. Because only when the user gets what he expects, only then, the user will trust the system because the reliability matches the accuracy of the system to fulfill the task.

At this point we started to draw our final conclusion and develop a theory around them.



Conclusions

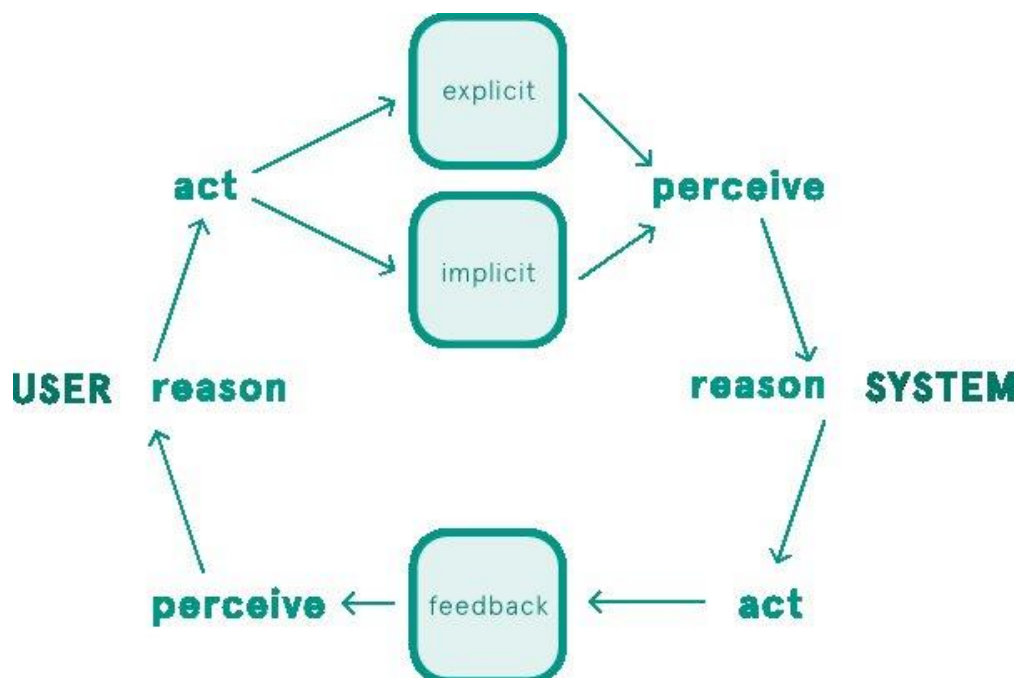
Mutual Intelligibility

In our opinion, the mutual intelligibility between user and system is one of the, if not *the* most important aspect of the interaction in an intelligent system. Within this interaction, both the user and the system give each other feedback in different ways.

You expect from intelligent systems they will understand your implicit feedback, but at the same time it does or might require your explicit input in order to learn. Some mistakes need to be fixed by the user manually because the system might not be smart enough, yet.

At the same time the system gives the user constantly feedback, but not every feedback is perceived properly by the user. This might be due to the fact that not all feedback can be noticed directly, just like with the temperature change with the Nest. It is something that happens overtime and is not that noticeable compared to a light turning on, off or changing colors. In cases like these it might be a very smart move by the system to give the user some extra feedback because the action wasn't enough feedback. This is also incorporated in our view that every action is feedback, but not all feedback is action and therefor does not have to be so.

In the end it is all about understanding each other and when looking at where things could go wrong we see almost identical mistakes on both sides: the system detects wrong user feedback or interprets it differently, and the system does not give correct feedback or feedback at all. At the same time the user might not give correct feedback or feedback at all themselves, and could interpret the feedback given by the system wrongly or simply not at all.



Conclusions

Pre-Knowledge and the Mental Model Cycle

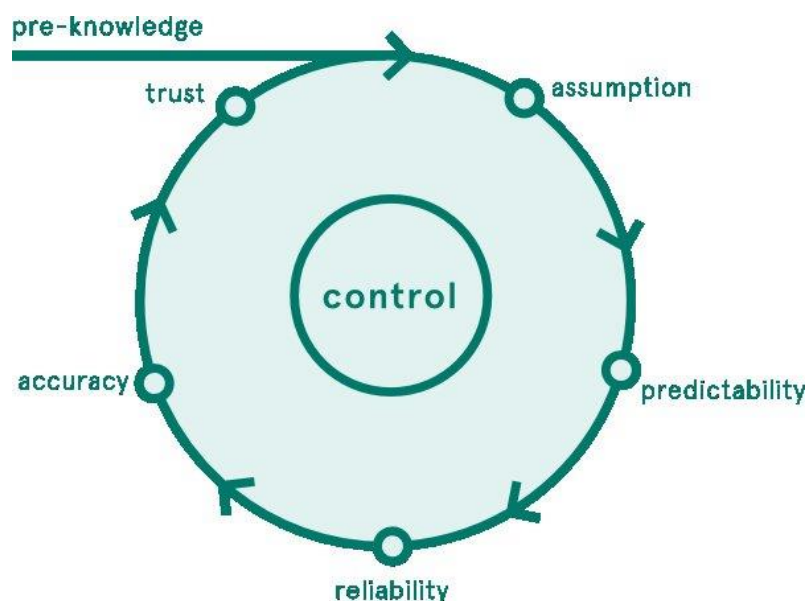
When users interact with a product, service or other artifact they will develop a mental model on how to interact with it. But many of the times we have some initial knowledge of a product, either by past experience of similar products or because of the affordance of the product. How does this pre-knowledge influence the mental model and how do we see the development of a mental model?

The interaction with a new product is largely based on something we already know. This pre-knowledge is based on things we've previously interacted with, it is based on things we already have a mental model for. As users we combine multiple mental models before we start using the product. These mental models form our pre-knowledge.

When interacting with the product, we start to understand which parts of the pre-knowledge is correct and which is not. From this we generate a new version of the pre-knowledge. It is essentially an adapted version of the mental models we had, turned into a specific mental model for this product. Only when we would interact with a completely new product, much like when we were very young children, we don't have pre-knowledge and we have to explore how something works. From this we then create a mental model.

The formation of this mental model is largely based on a desire for control. When interacting with the product we want to know what to do and what the product does. We have assumptions, which are based on our pre-knowledge, on what the product does, how it works and how we should interact with it. These assumptions allow us to predict how the product will behave, which then is either confirmed or debunked. This relation between our assumptions and the actual reliability create some sort of accuracy on the products behavior which we then use to either trust or distrust a product.

This cycle starts over again, except without the pre-knowledge. The pre-knowledge has already been made into a new mental model which we work with. During the cycle we adapt our assumptions (if needed) and this influences all the other parts of the cycle. Eventually we work towards maximized perceived control and depending on the product our mental model can be straight forward, or extremely complex.



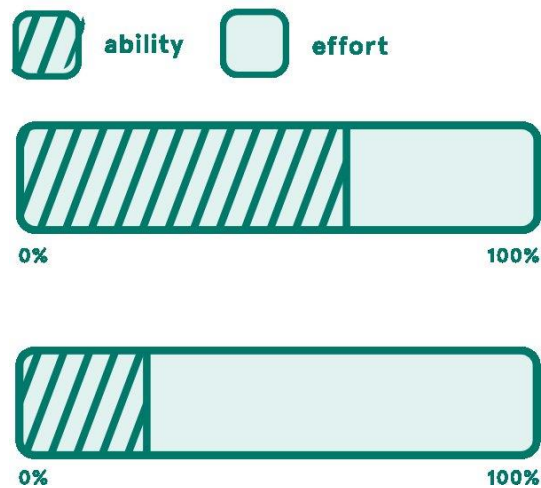
Conclusions

Expectation and Capability on Satisfaction

The development of a mental model is largely based on our ability to perform the tasks needed by the system, which then relates to the satisfaction of a system. Why would we want to use a system that is not satisfactory to our needs, may that be because of boringness in usage, the difficulty of the trust we have in the system.

In our view, the satisfaction is the relation between the effort needed and the capability of a user. When the effort is higher than the capability, satisfaction will be higher when successful. When capability is higher than effort, satisfaction will be less high or even low when successful.

At the same time, the mental model increases our ability to perform tasks within the system. It largely depends on whether the user wants to do the effort to develop more mental model complexity needed to attain specific abilities. This isn't necessary in every mental models, it all depends on what the user wants from a system.



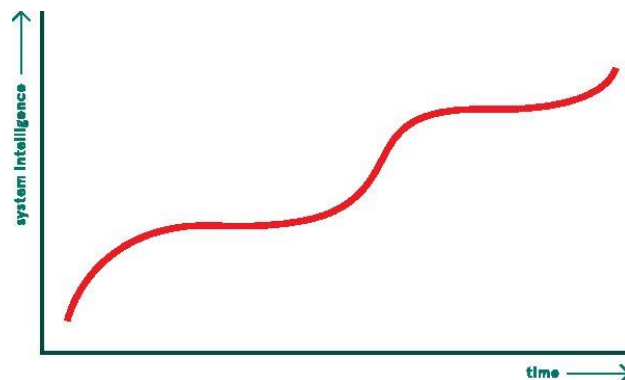
Conclusions

Intelligence versus Mental Model Complexity

At the start of the module we were given a specific question to answer: Is it really true that if we have intelligent systems, we can have less knowledge than if we have stupid system? Or do smart systems require different, more in-depth mental models with more knowledge? Or, does it depend based on different conditions?

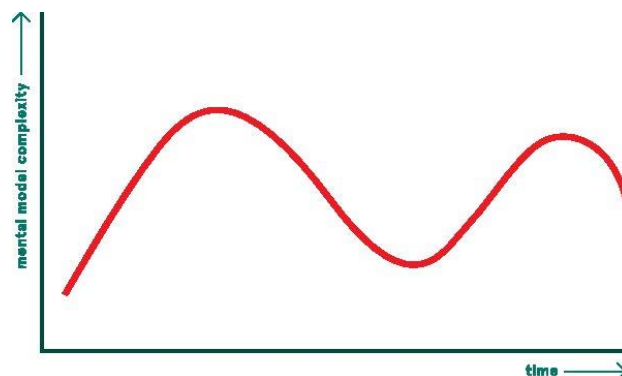
In this we want to make it clear there's a difference between the initial intelligence of the system and its adaptive intelligence. Initial intelligence is any knowledge an intelligent system has before being used by a consumer in order for it work correctly. All the things the system learns after that, is considered as part of its adaptive intelligence.

The system starts learning from the user, until the user is satisfied. The system will stop learning for a while until the user needs to system to learn again. However, at any point in time the system will never forget knowledge. Even when the system learns some forms of output are not desired by the user and it has to forget those ways of giving feedback, this is extra knowledge. It learns. The system becomes more intelligent and never less intelligent.



At the same time we believe that the mental model complexity is initially based on the pre-knowledge, which is a previous mental model, plus all the extra things the user needs to add to the mental model in order for the user to interact with the system as expected or required. At some points in the time, the user might forget some parts of this interaction, as was made clear during the telephone demo in the first lecture on how to make a call or how to set your ringtone. However, the user will always have the capability to relearn this based on the core of the mental model.

In essence we're saying the user's mental model complexity increases over time when more knowledge of a system is needed or wanted, but when this movement stops, the user will slowly loose some of its complexity only to be able to relearn it again on a later point in time.



However, not all systems require the same initial mental model complexity of the user. Through discussion and hypothesizing we concluded a specific pattern which relates the mental model complexity compared to the intelligence of the system. Keep in mind this is not over time.

When the system has a very low intelligence, the user needs some sort of mental model complexity. When the system has some more intelligence, the user needs more mental model complexity. However, at some point the system will reach a specific state of intelligence where the mental model complexity does not need to become more complex. This is because we believe that at this point, the system will start recognizing its own mistakes and the mistakes made by the user. The higher the intelligence, the better the capability of the system to recognize and fix these mistakes will become.

And at that point, the mental model complexity does not need to be so complex anymore. In a faraway future, where systems are in fact *that* smart, users don't need to understand the system in-depth anymore. The system will largely takeover much of the mental model complexity, as well as other human tasks.

Up until a certain point of intelligence, users need to become smarter with more system-specific knowledge. But when system intelligence thrives and development reaches higher and higher efficiencies, we *do not* need to be so smart and we can in fact become stupid ourselves.

