

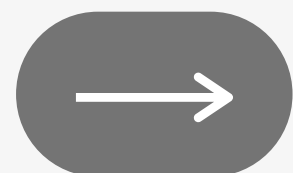
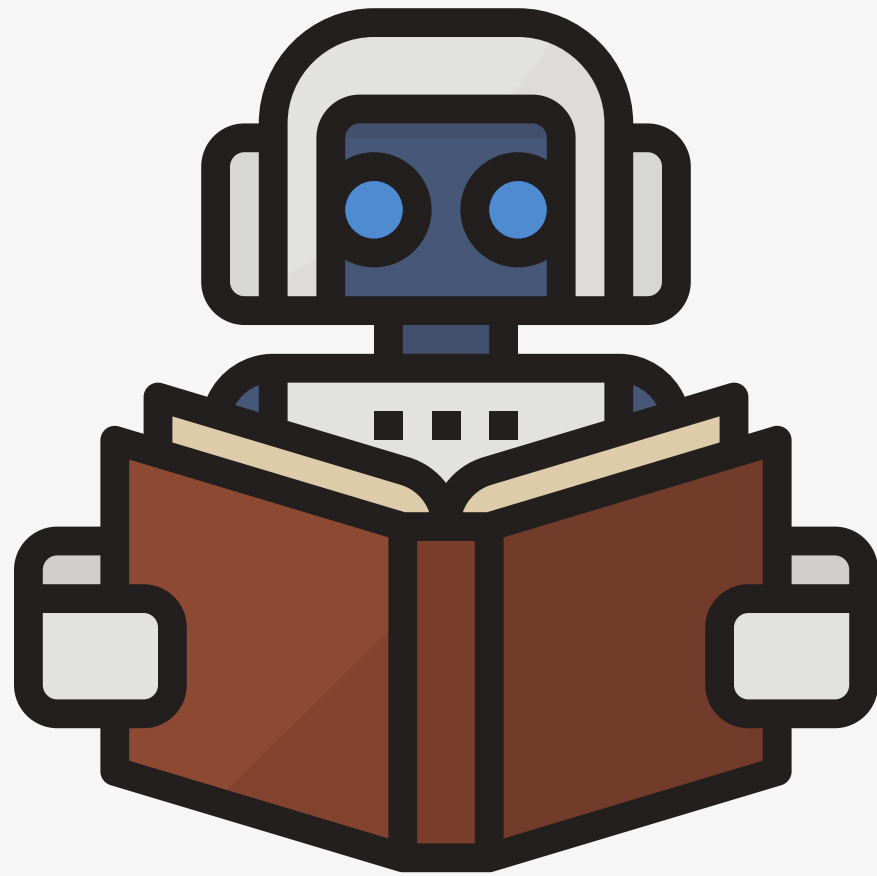


Machine Learning and IT Service Management

How can we use machine learning to free time of ITSM professionals?



Agenda



1

What are AI and Machine Learning?

2

How is Machine Learning used to drive Automation?

3

Main use cases within the ITSM space

SECTION 1

WHAT ARE AI AND MACHINE LEARNING?



AI

The terms 'Artificial Intelligence', 'Machine Learning' and 'Natural Language Processing' are all relevant in this context, but they are not the same thing.

Let's start with AI (Artificial Intelligence). As you can see on the right hand side, John McCarthy has defined the term AI for us. This may be a bit of an abstract definition, however as you continue reading through the rest of this section, things will start making more sense.

Now, there are many areas within the broad field of artificial intelligence, but one of them, the one that seems especially promising is Machine learning, which we'll go over in more detail now.

"It is the science and engineering of making intelligent machines, especially intelligent computer programs. It is related to the similar task of using computers to understand human intelligence, but AI does not have to confine itself to methods that are biologically observable."

- John McCarthy, 2004



Machine Learning

Now, so what exactly is Machine Learning?

IBM's definition can be found on the right hand side. It's very interesting, as this area in general likes to focus on trying to imitate and teach machines how to learn in a sort of similar way that humans do. Now, if you keep this in mind, the potential for self improvement on iteration will be something that will make how the variety of applications within the ITSM industry actually work, much clearer. We will talk about this in a bit more detail.

"Machine learning is a branch of AI and computer science which focuses on the use of data and algorithms to imitate the way that humans learn and gradually improve its accuracy"

- IBM Cloud Education



Natural Language Processing

There is also the concept of Natural Language Processing, which is another field of AI that focuses on understanding conceptual data. What this means is that we can use it to try to find meaning within unstructured data, like e-mails that can then be used for processing and within machine learning algorithms.

With a lot of data being unstructured in the context of the ITSM space, we will find that it will act as an enabler for us to work with this data.

Now, all of the things we have just mentioned help us drive automation, but it is important to recognise that not every piece of automation will be driven by AI. Naturally, a lot of tools have been automating various workflows for a long time, but it does not necessarily mean that they had anything to do with artificial intelligence and machine learning. This is the context that we will be focusing on and should act as the next stepping stone for the new generation of tools on the market or yet to enter the market.



So, Why Now?

1

More Data Available?

2

Moore's Law?

3

Better Algorithms?

All Three!



So, Why Now? (Continued)

So, where did the recent explosion and advancements in machine learning actually come from? It has been around for quite a while, in fact we were even aware of the possibilities since the last century, but it did not see as much success like it does now.

First of all, with recent tools storing more and more information, we have an abundance of data we can use to simply draw conclusions from. In addition, with a lot of applications being centralised and in the cloud, it is easier than ever to collect data about anything and everything.

Second of all, our computational capacity has increased immensely - CPUs have been getting faster and faster, year on year, meaning that the amount of computation is on a completely different level to what was possible in the past when machine learning was originally thought about. And, this is where Moore's law comes in - it is related to the number of transistors per silicone chip doubling roughly every couple of years. But the thing we are really interested in here, is the computational capacity this comes with.

Also, naturally, as the time went by there has been more research going in this field, which means that better algorithms across various domains were developed.

Now, this is certainly the time to use all of these resources to help to push industries forward, including the ITSM space.



SECTION 2

HOW IS MACHINE LEARNING USED TO DRIVE AUTOMATION?



Where is Machine Learning Actually Used?

Although this eBook is going to be generally focused on the applications within the ITSM space, and this is what the entirety of the next section is about, we are going to talk now about a variety of familiar examples in this part of the talk. Now, I hope that this will help to drive intuition and help people recognise where artificial intelligence is already used around us.

Let's say that you have finished working and want to relax for the night - now, many people will turn to Netflix, and start watching some TV shows/films. What many people don't know, is how Netflix is so good at getting us to do just that. You see, Netflix uses machine learning to make recommendations based on what you have already seen and grouping you with other people that liked and watched similar shows/films. By placing you in a specific group and recommending other shows/films that they think you will enjoy watching, they are able to influence you to spend more time using their service. Now, this may sound a bit off, but let's be honest, it does make you enjoy watching shows/films that Netflix recommends a lot more than you would otherwise.

Now, Netflix isn't the only company using AI driven recommendations. Another one is YouTube, where their models are, or at least in the past, have been tasked with recommending videos in such a way, to maximize the users' watch time. Meaning, when you are up at 2 am at night, wondering why you are watching more videos and why are they so interesting - it is because YouTube uses a range of algorithms to predict what you most likely want to watch, and what you will watch for the longest period of time.

When we are using services like Facebook or Google, the ads that we are seeing are more and more relevant to us. From the point of view of these companies, they use a bunch of data on us that they have stored, to try and determine what 'type' of person we are, and perhaps categorise us. Then, based on these different labels, you are likely going to see advertisements that will be best fit for you based on what other people with similar profiles have purchased in the past or what services they found interesting.

Now, it is not just necessarily driving profits that machine learning is used for. More and more, we see it applied across a variety of industries like healthcare. A recent example is a tool within the dermatology space launched by Google, which helps with diagnosing skin diseases.

So, clearly artificial intelligence is emerging across all of the sectors around us, making our lives more convenient and helping us access range of services that are relevant to us.

Great, this was the general introduction. Now let's move on to the next part, where we will cover two of the most common problem spaces within the machine learning domain - supervised learning, and unsupervised learning.

Don't get me wrong, there are many different areas and subcategories, like semi-supervised learning, reinforcement learning as well as deep learning, but we will focus on these two in this eBook. These two areas of machine learning are relevant and have strong examples that can be applied within the ITSM space.



Key Terms

Before we dig deeper, let's first discuss a few of the key terms that will pop up and will be important to understand:

Label

In supervised learning, 'label' refers to the result portion of the dataset. For example, if you had a set of data which included the information from a support ticket that was logged with by the end user, as well as how they were classified by technicians, and you wanted to automate the process of classification, the category would be the 'label', or the 'result' of this dataset.

Algorithm

This is just a procedure. You can think of it as a series of steps which explain how to achieve something. Imagine you were at your local library and wanted to walk back home. The directions, or the series of steps and turns you would take in order to get there would be the algorithm that describes how to get from point A to B. In machine learning, this is more specifically the series of steps run to create a model.

Model

This is the output of machine learning algorithms, which is made up of model data and algorithm predictions. It can comprise of any rules and data structures that will be used when running the model to make predictions. We will soon discuss an example of this, once we come to supervised learning.

Excellent, so with these key terms in mind, let's get onto supervised learning and unsupervised learning.



Supervised Learning

So, let's talk about supervised learning in more depth. We have mentioned it briefly when we came across the key terms, namely label, which is the result part of a data set for supervised learning.

Supervised learning focuses on learning from clean data, based on what sort of inputs, outputs, results, or labels we would expect. Clean data is used for training and to produce a model, and then usually another part of the dataset is used for validation. Because of this need for clean data, it can often be costly as we need experts within a given domain to label it, as well needing staff to clean it, meaning removing any noise and putting it in the correct format. Think of medical imaging, if we wanted to train a model to predict, based on x-ray scans, whether someone has a tumour, we would need to get a few doctors to go through a range of example pictures and label them, deciding whether or not a given picture was a good example.

Another problem with all of this is that we can overfit our model, which means that it won't be good at generalising and making predictions of things that are dissimilar to what it has seen in the training data set. This is where the split in data for the validation data set part comes in, so we can validate it, and make sure that the model performs to the desired accuracy, or whether we perhaps need to extend our training data set with a larger variety of better data.

To sum up, in supervised learning we focus on figuring out mappings between inputs and outputs with the two main problem domains within this area being classification, which deals with categorising and labelling, and regression, which deals with continuous prediction.



Supervised Learning (Continued)

Classification:

A good example of classification is when we take an incident or service request, and based on the pre-populated fields by the end user, we predict its category or perhaps we re-classify its ITIL type, which will help us route it to the correct team.

Regression:

Now, just to keep the examples aligned, a good example for regression is predicting what agent is likely going to spend the least time to resolve a ticket based on their past performance with similar tickets, and then routing the request to them. So, in this case regression would be used to decide which one of the available agents is likely to need the least amount of time for a given ticket.

Now, there are some drawbacks to supervised learning, apart from a lot of data being required and labelled, human errors can also introduce bias into the model. What this means is that if we had a technician that would misclassify certain tickets, and it wasn't picked up on, then later on, this data would be used to train our model and the model could carry the same bias that the agent did.



Unsupervised Learning

In contrast to supervised learning, the data that we use for training unsupervised learning models lack when it comes to labels. We do not say exactly what it is that we are looking for, but instead we allow it to explore and see what sort of conclusions it can draw out of data on its own.

Within this area, we focus more on looking for patterns within the data as the algorithms may not specifically know what it is looking for. Consequently, it may lead to surprising results.

A good sort of human analogy to use here, which would make things clearer would be to look at the image below. Think about the last time you had a chance to eat some skittles. Not sure how about you, but quite often people find themselves starting to separate them based on their colour, and start picking them out, and grouping them. And then naturally, eating them. The important bit here, is that it is likely that no one has shown us how to do it, or taught us. However, we look at them and we see that they carry specific characteristics, which in this case is colour, and as a result, proceed to cluster them based on this similarity that enables us to draw a conclusion that there are different subgroups amongst them.

What was just described relates to an application of unsupervised learning, specifically 'clustering'.



Clustering:

Clustering focuses on finding different subgroups within a dataset by looking at similarities and differences between different data points in a multi-dimensional space.

Now, multi-dimensional space may sound quite complex, but let's get back to our familiar example of a ticket. It is likely going to have a lot of information on it, categories, related assets, and so on. These can be individually compared across different tickets and although different tickets may have the same/similar categories, they may be very different with regards to what asset they relate to, or in fact, they may relate to the same asset.

We can perform various calculations to try to quantify this and predict which tickets are most likely related, and cluster them based on that and classify them into different subgroups.

There are many things we can do with this information, which is something we will cover shortly. However, to give you a rough idea, it includes possibilities of detecting repeatable service requests that are candidates for automation, as well as related incidents which may indicate that there is a larger problem.

This now concludes a very brief introduction to machine learning and more specifically to two of its subdomains. We have covered briefly supervised learning and unsupervised learning, with some mentions of how they can be applied within the ITSM domain, which we will now discuss in a lot more detail. We have also briefly discussed a range of applications within other industries and where they can be seen in our day to day lives.



SECTION 3

MAIN USE CASES WITHIN THE ITSM SPACE

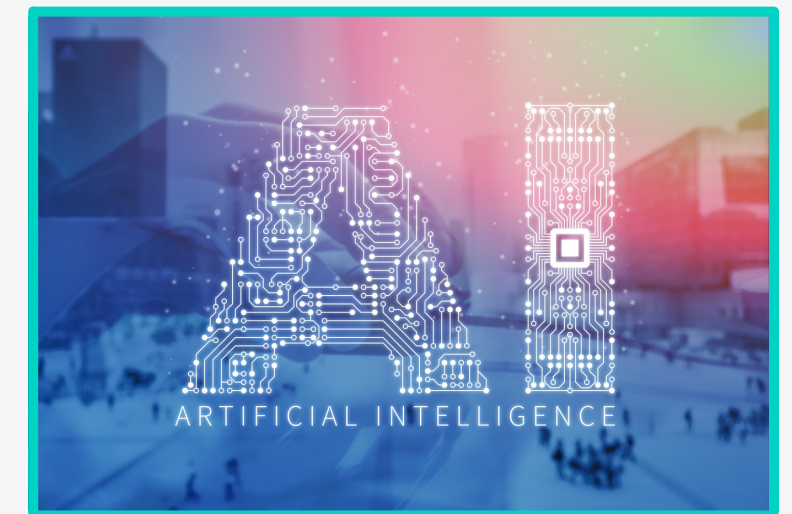


AI in ITSM

Now that we have discussed some of the basics of machine learning, we will now focus on how all of this fits into the ITSM industry, and where it is heading.

When it comes to ITSM, we want to drive the best user experience with the fastest resolution time possible for tickets. Things like end-user portals have been around for quite a while, trying to optimize the end-user experience and reduce the resolution time for tickets by providing knowledge base articles and reporting capabilities there, as well as freeing technicians' time to pursue a larger variety of tasks. However, as the industry grows, this is not sufficient as many service desks find themselves swamped with level 1 tickets that are very repeatable.

With all of this historical data stored, let's discuss how we can use it to bring the ITSM service to the next level, and leverage the new possibilities offered to us by machine learning, and more generally AI, to tackle these problems.



AI in ITSM

As mentioned in the introduction, we will be focusing on handling level 1 support, and asset lifecycle management. There will be a much larger range of use cases, but we will frame the various applications within the industry around these two aspects.

So, let's start with handling level one tickets. Now, with the emergence of knowledge base lookups and similar open and closed incidents, that a variety of systems can provide, we were successfully able to reduce the amount of time required to handle a significant amount of level one tickets based on historical data.

However, this is one of the areas where new applications of AI in the industry shine through and will hopefully make handling level 1 tickets a minor part of technicians' jobs, due to the new provided possibilities.



AI in ITSM

Let's start with one of the things just mentioned, knowledge base lookups. Many tools to date provide the capacity to search for knowledge base articles as a variety of tickets are logged. Now is the time to take it a step further. Rather than just providing suggested articles, we can train models to search for answers to the posted questions or tickets to improve this process further. This may sound really complex in terms of how we would approach this, but there are many specialists in the field who focus on natural language processing to help tackle the understanding of human languages to an extent that allows us to do such things.

One simple application and possible implementation on a toy case would be to use tensorflow, which allows us to look for answers to questions, based on available text. We could utilize it to provide suggested solutions based on the knowledge base articles accumulated by our organisations to reduce the resolution time and necessary involvement from ITSM professionals (For more example use cases, check out - <https://www.tensorflow.org/js/models>).

Now, this could be presented in the form of a pop-up when an end-user chooses to log a ticket, but it also goes very well with other emerging applications of artificial intelligence like virtual assistants and chatbots.

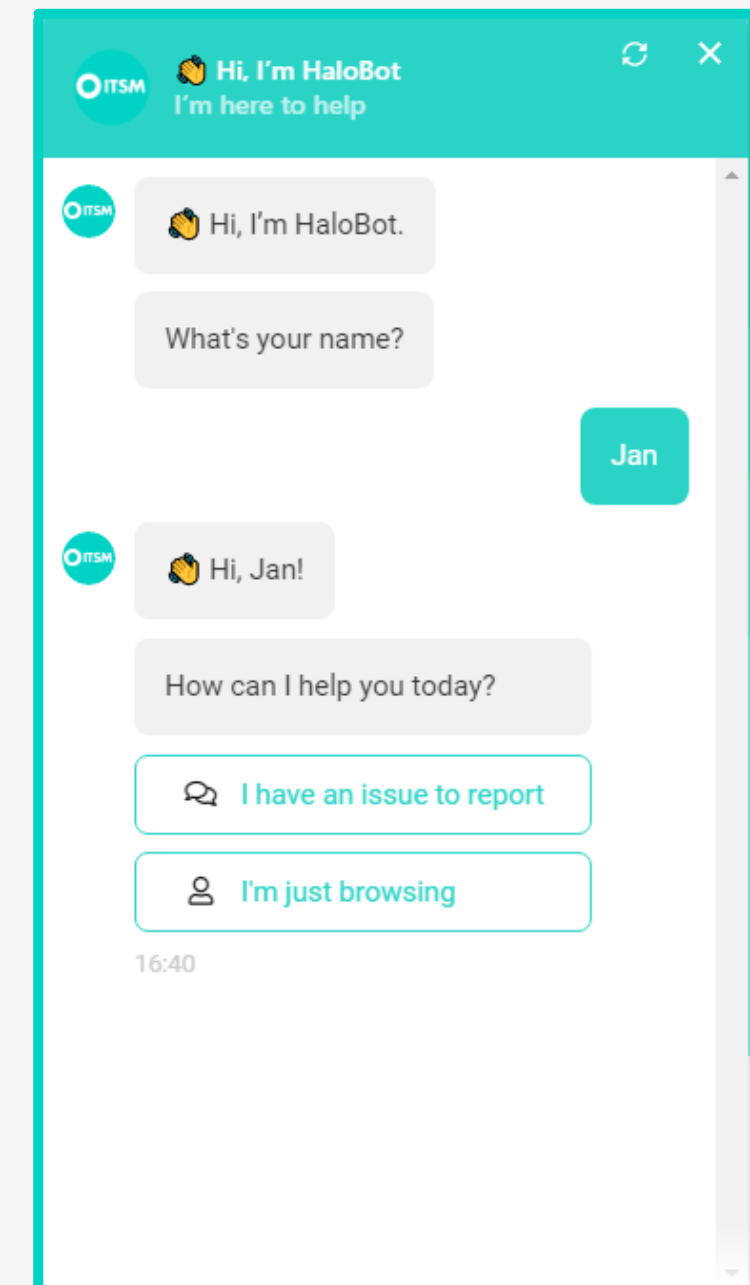


AI in ITSM (Virtual Assistants and Chatbots)

Whilst providing an interactive experience for end users, virtual assistants and chatbots can help them resolve their tickets on the spot. As our knowledge bases grow and as this technology improves with more improved data and algorithms, we will see that these will also become better, providing 24/7 support and only logging tickets that need human attention. Naturally, such support is cheaper, and with faster resolutions, end-user satisfaction goes up.

These can continue to improve by seeing which suggested solutions helped users and which didn't help, and then ranking these for the future uses, making the processes smoother. This is especially important as it also means that the tickets that are going to be logged for human attention will, hopefully with time, create new knowledge base entries, which will lead to higher automation and help fill in the organisations' knowledge gaps.

And for the tickets that get logged, what we can do is extract as much information as we possibly can from them to help categorize them correctly, as well as route them to the agents that have a history of resolving tickets within similar domains, in the smallest amount of time possible.



AI in ITSM (Smart Ticketing)

Now, to walk through a little bit of an example, imagine an end-user logs into an internal portal, they see a chatbot and decide to message saying that there's a problem and that they cannot access the internal network. The smart assistant informs the end-user about a range of possible solutions, and if none of them work, the smart assistant logs a ticket based on all of the collected information relating to the issue. All of the information from the ticket is processed, and based on this, it gets categorised and it gets assigned to a networking technician, who resolves the novel problem, and then creates a new article for it.

This should kind of give you a feel of how such a workflow works.

All of the automation around dealing with level one tickets we just discussed will certainly lead to:

- Improved end-user satisfaction
- Filled in knowledge base gaps across the organisation
- Reduced involvement of technicians
- Reduced support costs as a result by allowing ITSM professionals to pursue more meaningful and business driven tasks.



AI in ITSM – How much repetition is there?

Now, to quote Bertrand Lafforgue, "40% of calls to IT Service Desks are for recurring or similar cases". This is a lot of time that ITSM professionals could use to handle more important tasks.

Improved Interactions

In terms of support, when it comes to interactions technicians have with end-users, we can analyse any externally facing communications (or even internally if we wish) to advise on the tone of the message and how it could be interpreted by other people, therefore improving communication which is extremely important. This means that agents will soon be able to gain a better understanding into the way they explain a topic or provide an answer. Again, this is a very simple application and something that's really easy to implement. There are pretrained models that you can use for some toy play coming from TensorFlow as well.



Improved Asset Management Lifecycle

What can we do with all the data that we have been collecting?

Let's change things up a little bit and switch to the topic of assets. Now, if you have a modern ITSM system, you would have been collecting data about any tickets against your assets for quite a while now, the question is what can we actually do with all of this historical data?

Using all of this data, we can see if there are any recurring problems regarding particular technological assets or perhaps families of assets. This will help us to detect when certain assets may need replaced in order to ensure that their degraded performance does not restrict their usability and performance of any staff using them. We can also look at other trends, such as monitoring what software is typically installed on different assets and ensuring that any role that needs certain software will have a powerful enough machine in order to be able to run it.



Improved Asset Management Lifecycle

Problem Detection and Service Automation

Now, we have just touched on the topic of problems. Let's elaborate on this and bring back something that we briefly discussed when we were talking about clustering in unsupervised learning during the machine learning section.

As we have more and more tickets logged and created in our system, we can cluster them and try to look for any areas where there is a large amount of these, which would indicate that there is some sort of a relation.

Now, this helps us on two levels. The first one is if this large sample of tickets that are very similar get routed through to the agents, they are an excellent candidate for automation to be handled in the future by chat bots.

Now, this does not only apply for incidents, these things could be tickets along the lines of 'please, reset my password', which could be carried out by the virtual assistant amongst other things that would be identified during this process.

The second one is identifying any problems and seeing if any underlying root causes for a certain family of incidents can be proactively resolved. This is especially important as even if we have few tickets logged that are related, it will not necessarily be picked up on very quickly.

With use of machine learning, we will be able to tell whether there is a problem much faster and then proactively address it, resolve it, and communicate it to the end-user, if applicable. Alternatively, if they are related to a technological asset of some sort, let's say for a company that handles their own infrastructure rather than using one of the cloud providers, they should perhaps replace the faulty machine or server, or action it in one way or another.

Now again, by automating repeatable service requests, and resolving root causes for problems, we will be able to save the time of our staff, whilst improving end-user satisfaction through faster resolution times and hopefully lower incident rates.



Improved Asset Management Lifecycle

So, now we have had a brief discussion regarding the automation of level one ticketing and asset management lifecycle improvements, problem predictions, as well as various topics including virtual agents and smart ticketing.

There are many more use cases within the ITSM space, like smart searches with improved results shown to the searcher based on what other people searched for and what they did and did not find useful. In addition to this, resource requirement predictions based on historical data on service desk demand would allow to predict the required staffing.

Hopefully, you can clearly see how beneficial these technologies are for an organisation when it comes to meeting the ever increasing task volume, and freeing the time of ITSM professionals. Key benefits discussed include automating mundane, repeatable tasks, which will not only improve job satisfaction, but also allow them to pursue more business oriented tasks, while improving the end-user satisfaction.



Summary

First of all, we have briefly discussed AI and machine learning, as well as a couple of its branches. We have covered some applications within a range of industries, which should have hopefully given you a good understanding of the range of possibilities and cases in which we can apply it.

We have then discussed supervised learning, where the focus was on learning mappings between different inputs and outputs to try to predict the best answer. Think back to categorizing tickets based on the available data. Based on the training set and the available parameters we have used, we can find what given category fits best for a given ticket.

Then, we discussed unsupervised learning, which focused on looking for patterns within the data we passed on to the algorithm, but without us explicitly saying what it is that we are looking for. To go back to our original example, we spoke about clustering here, and how we could use it on tickets to try to predict which service requests could be automated, and later on in the ITSM part, we discussed how it could be achieved using virtual assistants, as well as how it could be used in proactive problem detection.

After this, we covered a variety of applications of AI and machine learning within the ITSM space and we discussed how these different features help save time for TSM professionals, and drive improvements across businesses. These included automated ticket routing and assignment, auto populating relevant information on the ticket from other bits of information available, chat bots providing level one support and interactive experience for users, improved management of assets and proactive management of problems, and knowledge management improvements among other things.

Now, as discussed, there are some drawbacks and difficulties with implementing these, ranging from potential high costs when it comes to development, and cleaning and labelling data. However, many tools on the market provide these capacities.

To quote Julie L. Mohr, "The potential for AI is significant, allowing the service provider to improve the customer experience, improve services in meaningful ways to the business, and shift the IT workforce from repetitive transactional work to innovative and creative work". As time goes by, computational capacity is going to improve, we are going to collect more data, we will have better algorithms, which will help to refine models even more, leading to continuous improvement in the industry and freeing the time of staff to pursue more meaningful tasks in a proactive way. Naturally, not only will it improve staff satisfaction, but also end-user experience due to lower resolution times, lower error rates with increased accuracy, and better decisions being made by organisations based on insights from the historical data.





Contact Us

Want to learn more about HaloITSM?

Phone Number

→ +44(0)1449 833 111 (UK)
+1 (619) 432-0470 (USA)
+61 (0) 388 205 182 (Australia)

Email Address

→ sales@imaginehalo.com

Website

→ www.haloitsm.com

HALOITSM

San Diego, USA | Melbourne, Australia | Stowmarket, UK