



XII Simpósio Brasileiro de Sistemas de Informação

**Tópicos em Sistemas de Informação:
Minicursos SBSI 2016**

Florianópolis/SC – 17 a 20 de maio de 2016

Sistemas de Informação na Era da Computação em Nuvem



XII Simpósio Brasileiro de Sistemas de Informação

De 17 a 20 de maio de 2016

Florianópolis – SC

Tópicos em Sistemas de Informação: Minicursos SBSI 2016

Sociedade Brasileira de Computação – SBC

Organizadores

Clodis Boscarioli

Ronaldo dos Santos Mello

Frank Augusto Siqueira

Patrícia Vilain

Realização

INE/UFSC – Departamento de Informática e Estatística/

Universidade Federal de Santa Catarina

Promoção

Sociedade Brasileira de Computação – SBC

Patrocínio Institucional

CAPES – Coordenação de Aperfeiçoamento de Pessoal de Nível Superior

CNPq - Conselho Nacional de Desenvolvimento Científico e Tecnológico

FAPESC - Fundação de Amparo à Pesquisa e Inovação do Estado de Santa Catarina

**Catálogo na fonte pela Biblioteca Universitária
da
Universidade Federal de Santa Catarina**

S612a Simpósio Brasileiro de Sistemas de Informação
(12. : 2016 : Florianópolis, SC)

Anais [do] XII Simpósio Brasileiro de Sistemas de Informação
[recurso eletrônico] / Tópicos em Sistemas de Informação: Minicursos
SBSI 2016 ; organizadores Clodis Boscarioli...[et al.] ; realização
Departamento de Informática e Estatística/ Universidade Federal de Santa
Catarina ; promoção Sociedade Brasileira de Computação (SBC). –
Florianópolis : UFSC/Departamento de Informática e Estatística, 2016.
1 e-book

Minicursos SBSI 2016: Tópicos em Sistemas de Informação

Disponível em: <http://sbsi2016.ufsc.br/anais/>

Evento realizado em Florianópolis de 17 a 20 de maio de 2016.

ISBN 978-85-7669-317-8

1. Sistemas de recuperação da informação – Congressos. 2. Tecnologia
– Serviços de informação – Congressos. 3. Internet na administração
pública – Congressos. I. Boscarioli, Clodis. II. Universidade Federal de
Santa Catarina. Departamento de Informática e Estatística. III. Sociedade
Brasileira de Computação. IV. Título.

CDU: 004.65

Prefácio

Dentre as atividades de Simpósio Brasileiro de Sistemas de Informação (SBSI) a discussão de temas atuais sobre pesquisa e ensino, e também sua relação com a indústria, é sempre oportunizada. Neste âmbito, os minicursos constituem-se para a comunidade em oportunidades de formação e atualização em determinados tópicos de pesquisa.

Na edição 2016 foram selecionados quatro minicursos dentre dezessete propostas submetidas por meio de uma chamada pública amplamente divulgada por meio de listas eletrônicas da SBC (Sociedade Brasileira de Computação). Todas as propostas receberam, no mínimo, três avaliações realizadas por um Comitê composto por quarenta e cinco professores doutores, que consideraram critérios como relevância para o evento, expectativa do público, atualidade e conteúdo. O SBSI 2016 inova também por lançar, pela primeira vez, o conteúdo dos minicursos no formato de livro, cujas quatro propostas selecionadas constituem os seus capítulos.

O Capítulo 1 apresenta o texto do minicurso intitulado **Análise de sentimentos utilizando técnicas de classificação multiclasse**, que visa introduzir o processo de análise de sentimentos utilizando técnicas de extração de características, executando as fases de pré-processamento textual, seleção de características, vetorização e, por fim, o aprendizado de máquina, a fim de inferir se uma opinião é positiva ou negativa (classificação binária) ou inferir um *rating* (classificação multiclasse).

O Capítulo 2 apresenta o texto do minicurso intitulado **Keep calm and visualize your data: minicurso de Visualização de Dados**, que apresenta uma abordagem prática que concilia a literatura em visualização de informação, e seus diversos recursos disponíveis, na construção de representações gráficas de dados que sejam facilitadoras para o processo decisório e a disseminação do conhecimento.

O Capítulo 3 apresenta o texto do minicurso intitulado **Modelos de Negócios Inovadores na Era da Computação em Nuvem**, no qual os autores apresentam modelos de negócios baseados em nuvem a fim de prover meios e conhecimento para a produção de um modelo de negócio inovador. Ainda, apresentam como e onde incubar e acelerar seu negócio, onde encontrar investidores, além de casos reais de modelos de negócio de sucesso e de fracassos, bem como o papel da computação em nuvem nesse processo.

O Capítulo 4 apresenta o texto do minicurso intitulado **Sistemas LBS, Internet das Coisas e Computação Vestível: Usando a Computação Sensível ao Contexto para Desenvolver as Aplicações do Séc. XXI**, que propicia uma visão geral sobre a Computação Móvel, Computação Sensível ao Contexto e Computação Ubíqua/Pervasiva, especialmente dos Sistemas Baseados em Localização (LBS), ilustrando suas potencialidades, problemas e desafios para que se tornem populares e atinjam, de fato, o usuário final de forma mais eficiente.

Acreditamos que este material poderá ser amplamente utilizado por professores de Sistemas de Informação, como base para suas aulas; por pesquisadores, para discussão de novas abordagens e como insumos para suas pesquisas atuais e futuras; e por profissionais, como auxílio em sua prática cotidiana. Esperamos que todos que tenham acesso ao conteúdo desses minicursos e façam dele um ótimo proveito!

Clodis Boscarioli (UNIOESTE) e Ronaldo dos Santos Mello (UFSC)
Organizadores da Trilha de Minicursos do SBSI 2016

Comitê de Programa

Adolfo Duran	UFBA
Avanilde Kemczinski	UDESC
Carla Merkle Westphall	UFSC
Carlos Alberto Vieira Campos	UNIRIO
Carlos Eduardo Santos Pires	UFCG
Claudia Cappelli	UNIRIO
Clodis Boscarioli	UNIOESTE
Clodoaldo Lima	USP
Cristiano Maciel	UFMT
Daniela Barreiro Claro	UFBA
Debora Paiva	UFMS
Denis Silveira	UFPE
Edmundo Spoto	UFG
Elvis Fusco	UNIVEM
Fatima Nunes	USP
Fernando Braz	IFC
Flavia Santoro	UNIRIO
Flavio Oquendo	IRISA/UBS
Flávio Soares Corrêa da Silva	USP
Geiza M. Hamazaki da Silva	UNIRIO
Geraldo Xexéo	UFRJ
Glauco Carneiro	UNIFACS
Guilherme Galante	UNIOESTE
Isabela Gasparini	UDESC
João Porto de Albuquerque	USP
José Maria David	UFJF
Juliano Lopes de Oliveira	UFG
Lais Salvador	UFBA
Leonardo Azevedo	IBM/UNIRIO
Leticia Mara Peres	UFPR
Luciana Zaina	UFSCar
Luciano Digiampietri	USP
Lucineia Heloisa Thom	UFRGS
Marcos Chaim	USP
Maria Istela Cagnin	UFMS
Márcio Barros	UNIRIO
Morganna Diniz	UNIRIO
Patricia Vilain	UFSC
Regina Braga	UFJF
Ricardo Choren	IME
Rita Suzana Pitangueira Maciel	UFBA
Roberto Pereira	UFPR
Rodolfo Resende	UFMG
Ronaldo Mello	UFSC
Tania Tait	UEM

Sumário

Análise de Sentimentos Utilizando Técnicas de Classificação Multiclasse.....	1
<i>Alexandre de Castro Lunardi, José Viterbo, Flávia Cristina Bernardini (Universidade Federal Fluminense - UFF)</i>	
<i>Keep Calm and Visualize Your Data</i>	31
<i>Fernanda C. Ribeiro, Bárbara P. Caetano (Universidade Federal do Rio de Janeiro - UFRJ), Melise M. V. de Paula, Guilherme X. Ferreira, Rafael S. de Oliveira (Universidade Federal de Itajubá - UNIFEI)</i>	
Modelos de Negócios Inovadores na Era da Computação em Nuvem	53
<i>Ricardo Batista Rodrigues, Carlo M. R. Silva, Vinicius C. Garcia, Silvio R. L. Meira (Universidade Federal de Pernambuco - UFPE)</i>	
Sistemas LBS, Internet das Coisas e Computação Vestível: Usando a Computação Sensível ao Contexto para Desenvolver as Aplicações do Séc. XXI	72
<i>Marcio Pereira de Sá (Faculdade Sul-Americana - FASAM)</i>	

Capítulo

1

Análise de Sentimentos Utilizando Técnicas de Classificação Multiclasse

Alexandre de Castro Lunardi, José Viterbo, Flávia Cristina Bernardini

Abstract

After the advent of Web 2.0, finding reviews on products, businesses, services, organizations and many other areas on the web became really easy. These opinions can be found on social networks, blogs and specialized e-commerce sites that provide tools so that a user can evaluate an item. These reviews can be useful in recommender systems, stating whether a product is suitable or not via the existing web assessments. This type of analysis is known as binary classification. Beyond binary classification, other forms may include scales evaluation, problem known as multiclass classification. An example of this type of classification is ratings inference (multiclass classification). This work intends to introduce the feelings analysis process, using feature extraction techniques by performing the steps of textual preprocessing, feature selection, vectorization, and finally, machine learning, in order to infer whether an opinion is positive or negative (binary classification) or to infer a rating.

Resumo

Com o advento da Web 2.0, encontrar opiniões sobre produtos, negócios, serviços, organizações e sobre tantos outros domínios é algo comum na web. Essas opiniões podem ser encontradas em redes sociais, blogs especializados e em sites e-commerce que disponibilizam ferramentas para que um usuário possa avaliar um item. Essas opiniões podem ser de grande utilidade em sistemas de recomendação, informando se um produto é recomendado ou não por meio das avaliações existentes na web. Este tipo de análise é baseada em técnicas de classificação binária. Além da classificação binária, outras formas podem incluir a avaliação de escalas, problema conhecido como classificação multiclasse. Um exemplo deste tipo de classificação é a inferência de ratings (classificação multiclasse). Esse estudo tem o intuito de introduzir o processo de análise de sentimentos utilizando técnicas de extração de características, executando as fases de pré-processamento textual, seleção de características,

vetorização e, por fim, o aprendizado de máquina, a fim de inferir se uma opinião é positiva ou negativa (classificação binária) ou inferir um rating.

1.1 Introdução

Opiniões sempre foram úteis no que diz respeito à tomada de decisões dos seres humanos [CAMBRIA et al. 2013]. Nossas escolhas sempre foram, em certo grau, dependentes das opiniões e conselhos de outras pessoas [LIU 2012]. Além disso, é de grande importância para empresas conhecer o sentimento das pessoas em relação a um produto ou serviço, o que permite realizar previsões de mercado ou oferecer recomendações aos consumidores [TURNER 2002], tornando essas empresas mais próximas de seu público-alvo.

Com o advento da Web 2.0, é cada vez mais fácil encontrar opiniões valiosas relacionadas a produtos, serviços, organizações, indivíduos, eventos e vários outros domínios. Isso se deve ao crescente uso de redes sociais, blogs e, principalmente, ferramentas que permitem aos usuários deixar registrado seus comentários sobre algum produto ou serviço em sites de comércio eletrônico. Essa crescente disponibilização de dados é também conhecida como “web social” [CAMBRIA et al. 2013]. Isso pode ser notado em sites de comércio eletrônico como o *Booking.com*¹ e a *Amazon*², nos quais os clientes podem deixar seus comentários, revelando suas opiniões a respeito do produto ou serviço oferecido.

Com essa grande quantidade de informação disponível na Internet, analisar todo o conjunto de opiniões encontradas se tornou uma tarefa inviável para o ser humano. Com isso, capturar e processar de forma adequada essas informações por meio de técnicas computacionais – a chamada mineração de opiniões ou análise de sentimentos [CAMBRIA et al. 2013] – se torna fundamental para permitir a identificação do real interesse do público sobre algum item. A comunidade científica vem, dessa forma, desenvolvendo ferramentas que visam auxiliar na recuperação e tratamento de opiniões – ou avaliações – sobre produtos e serviços, disponíveis na web social. Dessa forma, pesquisas sobre mineração de opiniões e/ou análise de sentimentos são uma das áreas mais ativas e desafiantes, abordadas principalmente na área de Processamento de Linguagem Natural (NLP).

Análise de sentimentos ou *mineração de opiniões* são os principais termos empregados para descrever a análise automática de textos subjetivos, isto é, textos que contêm não apenas fatos ou explicações técnicas sobre algo, mas alguma opinião a respeito de um item. A partir dessa análise, é possível identificar aspectos distintos de um item, por exemplo, a localização ou a limpeza de um hotel, a durabilidade ou a facilidade de uso de um eletrodoméstico. É possível realizar também uma análise agregada das avaliações sobre um determinado item, identificando o sentimento geral em relação a esse item. Com isso, poderemos saber, por exemplo, se um hotel é recomendado pelos consumidores que ali se hospedaram, de acordo com o conjunto de opiniões emitidas. Ou seja, considerando-se apenas uma escala binária, os aspectos específicos de um produto ou serviço podem ser classificados como bons ou ruins. Por exemplo, pode-se avaliar se a câmera de um celular é recomendável ou não, ou se a localização de um hotel é boa ou

¹ <http://www.booking.com/>

² <http://www.amazon.com/>

não. A partir da análise de diversos aspectos, chega-se a uma conclusão sobre o sentimento final sobre um item. Por exemplo, após avaliar atributos de um celular como a câmera, facilidade de uso, preço e vários outros aspectos, pode-se chegar a uma conclusão final sobre a recomendação do celular.

Considerando um texto subjetivo que representa a avaliação de um usuário sobre um determinado item, este pode ser classificado utilizando-se técnicas binárias ou multiclasse. Na classificação binária, o objetivo é rotular essa avaliação como positiva ou negativa (boa ou ruim), com relação ao sentimento expressado pelo usuário. Um exemplo de ferramenta criada para analisar o sentimento de uma opinião é a *sentiment140*³, proposta por [GO; BHAYANI; HUANG 2009]. Nesse site é possível verificar o sentimento em relação a uma entidade (empresa, produto, serviço etc) utilizando *tweets* sobre a entidade em análise. Essa ferramenta seleciona os *tweets* de acordo com a palavra-chave informada pelo usuário e classifica os *tweets* encontrados como positivos ou negativos. Além disso, é apresentado um gráfico com a porcentagem total de *tweets* positivos e negativos. Outros sites também exemplificam o uso da análise de sentimentos, como o *NLTK Text Classification*⁴, no qual o usuário digita um texto sobre algo e o sistema determina se é uma opinião (texto subjetivo) e, caso positivo, se a polaridade desta é positiva ou negativa. Outro exemplo é o site *Skyttle*⁵, no qual a opinião informada também é classificada como boa ou ruim e, além disso, as frases com sentimento bom são marcadas em verde e as frases com sentimento ruim são marcadas em vermelho.

A classificação ou análise multiclasse analisa uma avaliação considerando escalas com mais de dois “valores de sentimento”, como por exemplo, “bom”, “neutro” ou “ruim”. Em diversos cenários, os produtos ou serviços são classificados em escalas de valores múltiplos. Pode-se citar como exemplo o site *Booking*, no qual os hotéis são classificados com notas que variam de 0 a 10. Assim sendo, uma forma típica de análise multiclasse é o Problema de Inferência de *Rating* (*Rating-inference Problem- RIP*), baseada em escalas de *rating* que tipicamente variam de 1 a 5 estrelas [PAN e LEE 2005]. Em alguns casos, essa escala pode ser analisada como 4 classes, na qual a classe 3 (neutra) é desconsiderada, ou como 3 classes, nas quais as classes 1 e 2 são unidas, assim como as classes 4 e 5.

Escalas baseadas em *ratings* estão presentes em larga escala em ferramentas de avaliação disponíveis em serviços como *Amazon*TM e *Netflix*^{TM6}, e projetos como o *GroupLens*^{TM7} [KONSTAN *et al.* 1997], com avaliações utilizando opiniões rotuladas em uma escala 5-*ratings*. A importância da avaliação correta pode ser comprovada de acordo com a pesquisa do site *ComScore*⁸, que mostra que os consumidores têm maior disposição para gastar entre 20% e 99% a mais em serviços que tenham uma classificação excelente (5 estrelas) do que um serviço classificado com 4 estrelas (Bom). Para o domínio de hotéis, esse percentual é de 38%. Técnicas de classificação binária não permitiriam que essa divisão (4 e 5 estrelas) fosse identificada automaticamente em trabalhos em análise de sentimentos. Além disso, considerando a grande utilização de várias classes na análise de textos subjetivos, é de grande interesse analisar não somente a polaridade de um item,

³ <http://www.sentiment140.com>

⁴ <http://text-processing.com/demo/sentiment/>

⁵ <http://www.skyttle.com/demoin>

⁶ <http://www.netflix.com>

⁷ <http://grouplens.org>

⁸ <http://comscore.com>

mas também avaliar os graus de positividade e negatividade por meio de *ratings* numéricos, baseados, por exemplo, na escala de Likert [LIKERT 1932], variando de 1 a 5 estrelas.

Embora seja uma forma de classificação essencial devido ao grande uso de *ratings* e de grande importância para a comunidade e para empresas, o número de trabalhos disponíveis em análise de sentimentos multiclasse é muito inferior se comparado aos trabalhos com classificação binária. Mesmo em cenários tipicamente de escalas múltiplas, grande parte das análises de opiniões consideram apenas duas classes principais – agrupadas em recomendado ou não recomendado –, em que, em uma escala de 1 a 5 estrelas, por exemplo, 4 ou 5 estrelas são consideradas recomendáveis e 1, 2 ou 3 estrelas não são recomendáveis.

Um possível emprego para as técnicas de análise multiclasse, seria permitir a classificação automática de comentários de usuários em sites de produtos ou serviços, diminuindo o chamado efeito manada (*herding effect*) [WANG e WANG 2014]. Esse efeito acontece na avaliação direta realizada pelos usuários quando estes se deixam influenciar pela avaliação da maioria. Por exemplo, um usuário pode ter achado um celular bom (4 estrelas), mas caso a maioria dos outros usuários tenham considerado excelente (5 estrelas), existe a possibilidade de que ele avalie o celular com base na média dos outros usuários e não no que o celular representou para ele. Além disso, a classificação automática de comentários baseada em análise multiclasse poderia simplesmente evitar os erros da classificação realizada pelo usuário, e casos em que um comentário não condiz o número de estrelas atribuídos, poderiam ser evitados. Mesmo com a divisão com *ratings* 4 e 5 consideradas recomendadas, de acordo com a pesquisa feita pelo site *PracticalECommerce*⁹ a inferência de *ratings* seria de grande utilidade tendo em vista que uma estrela a mais ou a menos pode fazer a diferença no momento da compra de um item.

Dessa forma, esse capítulo tem como principal objetivo a apresentação de técnicas de análise de sentimentos baseada em classificação multiclasse. Primeiramente apresentamos os principais conceitos relacionados ao tema. Em seguida, apresentamos as técnicas de extração de características que possibilitam uma boa representação de opiniões na forma de vetores de características. Em seguida, explicamos os diversos modelos e algoritmos de classificação que podem ser utilizados e discutimos métricas para avaliar de forma adequada o desempenho desses algoritmos. Além disso, apresentamos também uma discussão sobre os trabalhos que utilizam aprendizado de máquina na área de análise de sentimentos.

1.2 Definições

Segundo [LIU 2012], a *análise de sentimentos* ou *mineração de opiniões* é o campo de estudo que analisa as atitudes, emoções, sentimentos e as opiniões das pessoas em relação a entidades - como produtos, serviços, organizações, eventos, tópicos - e os atributos dessas entidades. Ela é um campo desafiador na área de Processamento de Linguagem Natural já que trata várias questões de PLN como a tratamento de negação e retirada de palavras-chave e pode cobrir muitos problemas, desde a classificação em relação à polaridade de uma opinião até o processo de sumarização do sentimento geral sobre algo.

⁹ <http://www.practicalecommerce.com/articles/93017-Study-5-Star-Reviews-Not-Necessarily-Helpful>

Seja um texto d , a tarefa inicial na mineração de opiniões consiste em determinar se d é subjetivo, ou seja, expressa um sentimento. Seja tal texto considerado subjetivo, formalmente, uma opinião pode ser representada como uma 5-tupla [LIU 2012] $O = (e, a, s, h, t)$, na qual:

- e é o nome da entidade ou objeto ao qual uma opinião se refere;
- a é o atributo específico da entidade;
- s é o sentimento do autor da opinião em relação a um atributo ou entidade;
- h é o autor da opinião, e;
- t é a data na qual a opinião foi criada.

Como exemplo, temos a seguinte opinião sobre o hotel H retirada do site TripAdvisor¹:

User: DesDeeMona (h)

Title: A magnicent building of fading grandeur, redolent of earlier times

Rating: 4

Date: April 28, 2015 (t)

Review: "Ground floor lobbies and suites with art deco design are impressive. Rooms (a) are spacious (s) with high ceilings and plenty of room in the en suite shower. Yes, lifts are a little slow, the paint is peeling, the plaster cracking, there are stains on the carpet - but hey, everything works, the sheets are well laundered and beds are comfortable".

Nesse exemplo, pode-se observar que o hotel é a entidade e um dos atributos são os quartos, destacados com a letra (a) no documento. Eles são classificados pelo autor como “espaçosos”, como demarcado acima pela letra (s) . Para [DAVE et al. 2003], a tarefa ideal na análise de sentimentos deveria processar um conjunto de opiniões sobre certa entidade, gerando uma lista de atributos para a mesma e agregar opiniões sobre os atributos da entidade. Entretanto, outros autores consideram apenas a entidade da opinião e o sentimento final, como feito em [PANG; LEE; VAITHYANATHAN 2002].

De forma resumida, a taxonomia de análise de sentimentos está presente na Figura 1. A ideia básica para o problema de análise de sentimentos é que um usuário emita uma opinião, também chamada de avaliação ou revisão. Essa avaliação pode ser sobre uma entidade ou item (como um hotel, por exemplo) ou pode ser relativa a um aspecto ou atributo específico de um item (a localização do hotel). Esse sentimento geralmente pode ser classificado em duas ou mais classes. A forma mais comum de classificação é a classificação binária, que diz se uma opinião é positiva ou negativa. Além disso, outras formas de classificação merecem destaque como a classificação por meio de *ratings* (presentes no site da *Amazon*) ou por meio de notas (*Booking.com*).

1.3 A Análise de Sentimentos e o Aprendizado de Máquina

Um dos primeiros trabalhos a analisar o sentimento das pessoas através de dados da web foi discutido em [DAS e CHEN 2001], e utilizou o termo *extração de sentimento* para capturar a influência da opinião de indivíduos no domínio de finanças. Já Pang et al. [PANG et al. 2002], utilizam o termo *classificação de sentimentos* para avaliar documentos considerando o sentimento geral de uma opinião, classificando-as como positivas ou negativas. Outro trabalho inicial é o de Turney [TURNERY 2002] que visa classificar opiniões como *recomendadas* ou *não recomendadas* (em inglês, *thumbs up* e *thumbs down*). Apenas em Nasukawa e Yi [NASUKAWA e YI 2003] o termo *análise de*

sentimentos é empregado, e assim como em [PANG; LEE; VAITHYANATHAN 2002], os autores introduzem uma pesquisa para classificar uma opinião como positiva ou negativa.

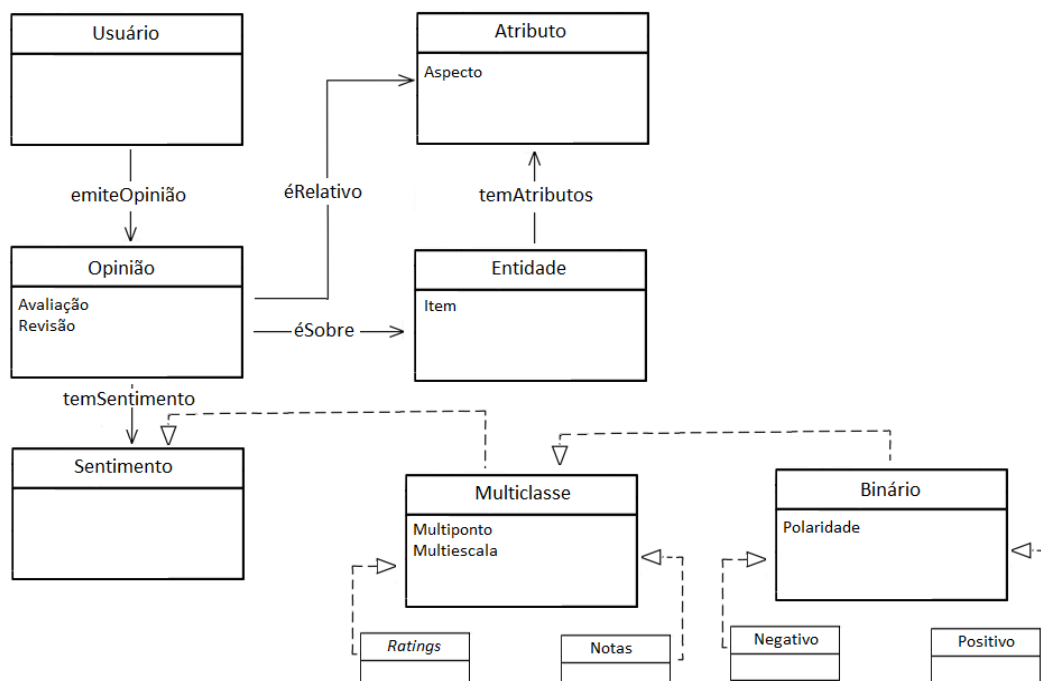


Figura 1. Ontologia com a análise de sentimentos

De acordo com Cambria *et al.* [CAMBRIA *et al.* 2013], a mineração de opiniões pode ser agrupada em quatro campos, na qual a análise pode ser realizada por meio de:

- **Palavras-chave e afinidade léxica:** classifica o texto de acordo com a presença de palavras sem sentido ambíguo, tais como “feliz”, “triste” e “medo”. Além de detectar palavras óbvias, também atribui a outras palavras uma relação de afinidade com um sentimento, seja ele bom ou ruim. Um exemplo de aplicação é o *SentiWordNet*¹⁰ 3.0 [BACCIANELLA; ESULI; SEBASTIANI 2010], um recurso léxico criado a fim de orientar aplicações em mineração de opiniões.
- **Aprendizado de máquina:** utiliza modelos de aprendizado de máquina, como Naive Bayes e Support Vector Machine (SVM), para classificar um texto. Nesse caso, o sistema, além de aprender a importância de uma palavra-chave óbvia, considera outras palavras que podem ser fundamentais, além da possibilidade de analisar a frequência ou a pontuação de um texto.
- **Orientação semântica:** esses métodos calculam a orientação semântica (por exemplo, para o problema binário saber a polaridade da palavra) de uma palavra baseada na coocorrência da mesma com palavras que possuem a mesma orientação. O principal trabalho que propõe um método que calcule essa orientação semântica é o algoritmo proposto por Turney, 2002 [TURNERY 2002]. O algoritmo Pointwise Mutual Information and Information Retrieval (PMI-IR) é utilizado a fim de medir a similaridade de pares de palavras ou frases. A orientação é calculada pela comparação da similaridade de uma palavra em relação aos sentimentos positivo e negativo.

¹⁰ <http://www.sentiwordnet.isti.cnr.it>

- Baseado em conceitos: usam ontologias ou redes de palavras-chave para realizar a análise textual. Podem analisar expressões que não possuem uma emoção explícita, mas estão relacionadas a um sentimento implicitamente. No trabalho realizado por Kontopoulos et al [KONTOPOULOS *et al.* 2013], é proposto o uso de ontologias a fim de melhorar o desempenho da análise de sentimentos no *Twitter*™.

Como pode ser notado, existem várias técnicas de análise de sentimentos, entretanto, o foco desse estudo está na utilização de modelos de aprendizado de máquina, juntamente com técnicas de extração de características, a fim de treinar e classificar um conjunto de opiniões, de acordo com o esquema exibido pela Figura 2.

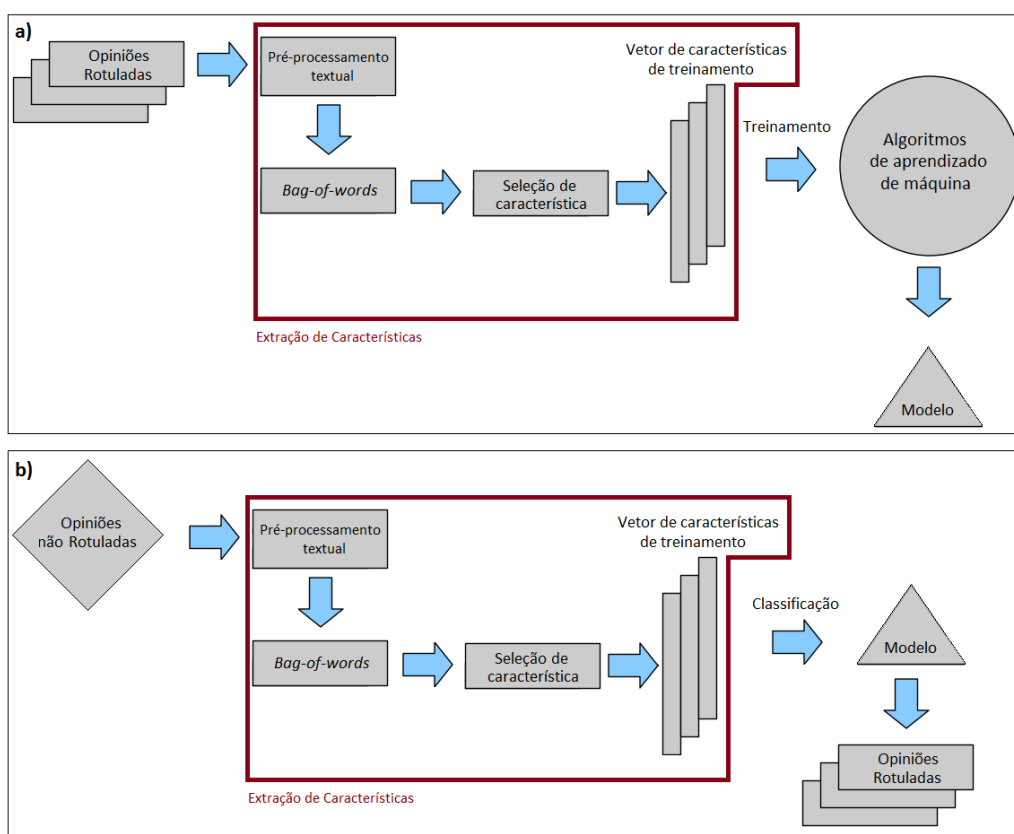


Figura 2. Processo de análise textual com extração de características e aprendizado de máquina. (a) Processo de treinamento. (b) Processo de classificação

Na parte a), o processo de extração de características e o treinamento dos algoritmos de aprendizado são descritos. Após a seleção de uma base de dados com opiniões previamente rotuladas, a fase de extração de características é dividida em quatro etapas. A primeira etapa, de pré-processamento textual, consiste na retirada de caracteres especiais, *stopwords* e tratamento da negação. A segunda etapa, *bag-of-words*, transforma cada opinião da base de dados em um conjunto de unigramas e bigramas. A terceira etapa é a fase de seleção de características que consiste na escolha dos melhores n-gramas para o treinamento dos algoritmos de classificação. A última etapa é a de vetorização que transforma a base de dados em documentos que são mais facilmente compreendidos pelos algoritmos de aprendizado. Por fim, esses algoritmos criam modelos de classificação que podem ser utilizados para categorizar opiniões sem rótulos.

A parte b) representa o modelo de classificação de novas instâncias. As opiniões não rotuladas selecionadas passam pelo mesmo processo de extração de características descrito na parte a). Após a extração de características, as opiniões são classificadas através de um modelo criado na parte a).

As etapas de pré-processamento, *bag-of-words*, seleção de características e vetorização de características para treinamento estão descritas na Seção 1.4. Essas técnicas estão presentes em alguns dos principais trabalhos que tem como objetivo a análise de sentimentos por meio do aprendizado de máquina. Na Seção 1.5, os principais algoritmos de classificação utilizados em análise de sentimentos são apresentados. Além disso, as medidas avaliativas são descritas na Seção 1.6.

Embora existam centenas de técnicas e métodos de análise de sentimentos presentes na literatura, como pode ser notado na Seção 1.7, que discute os trabalhos relacionados, alguns passos comuns e técnicas bem utilizadas foram selecionados baseado na importância dos trabalhos e no bom desempenho dos métodos existentes.

1.4 Extração de Características

Seja um conjunto D de opiniões selecionadas, algumas fases são essenciais na extração de características de textos que contêm opiniões. Embora as Técnicas de Extração de Características (TEC's) descritas a seguir não agreguem todas as formas de análise disponíveis, por meio destes passos é possível configurar um bom documento que possa ser compreendido por um algoritmo de aprendizado de máquina, cuja análise é o foco deste trabalho.

1.4.1 Pré-Processamento Textual

O primeiro passo para a construção de um documento compreensível para os algoritmos de classificação é selecionar uma base de dados com textos avaliativos, isto é, textos que possuam um sentimento em relação a um item.

Com as opiniões a serem analisadas devidamente selecionadas, o próximo passo do pré-processamento é a *tokenization*, que consiste na retirada de caracteres como vírgulas, acentos e pontuações. Em alguns trabalhos, alguns caracteres, como pontos de exclamação ou *emoticons* podem ser utilizados como característica de treinamento [GO; BHAYANI; HUANG 2009] ou como forma de seleção de opiniões para a criação de uma base de dados [PAK e PAROUBEK 2010]. Em casos nos quais as opiniões são extraídas diretamente de páginas web, a retirada de *tags* em HTML também deve ser realizada como feito em [BEINEKE *et al.* 2004] e [KANG; YOO; HAN 2012]. Alguns exemplos de caracteres especiais estão presentes na Tabela 1.

Tabela 1. Exemplos de caracteres especiais

Descrição	Token
Acentos	´ ~ ^
Pontuação	‘ ’ , . ; : ? !
Especiais	@ # * () &
Emoticons	:) ;) :D : (; (
HTML	 <p>

Com a retirada desses caracteres especiais das opiniões, o próximo passo é o da normalização textual. Nesta etapa, estão incluídas a retirada de radicais, retirada de letras

repetidas em algumas palavras e a correção ortográfica. A etapa de correção ortográfica pode ser notada em trabalhos como [KOULOUMPIS; WILSON; MOORE 2011]. Embora seja indiscutível a necessidade desse passo, poucos trabalhos que tem o foco na classificação de sentimentos citam a normalização textual na fase de extração de características. Esses passos podem ser melhor estudadas em livros como [MANNING e RAGHAVAN 2009] que, além de apresentarem uma boa introdução sobre recuperação de informação e o pré-processamento textual, mostram a utilização de algoritmos de aprendizado para a classificação textual.

Outro passo importante na parte de tratamento das opiniões é a retirada de palavras consideradas com pouco ou nenhum sentimento, as chamadas *stopwords*¹¹. O objetivo é diminuir a quantidade de palavras que possam ser usadas no treinamento, retirando palavras que pouco influenciam na determinação do sentimento final de um texto.

Outro passo importante é o tratamento de opiniões com palavras que expressam negação [PANG e LEE 2008]. Desta forma, frases como “*This is not bad*” ou “*That is not good*” tem seu sentimento invertido pelo *token* “*not*”. A fim de tratar esse problema, palavras que tem como precedentes os modificadores *no*, *not* ou *nothing* podem ser transformadas em uma única palavra. Como exemplo, “*not good*” é representado pelo *token* “*not_good*” que é similar ao *token* “*bad*”.

1.4.2 N-Gramas – *Bag of Words*

Com as opiniões normalizadas, cada palavra de uma opinião corresponde a um unigrama, como pode ser observado no trabalho de Pang *et al.* [PANG; LEE; VAITHYANATHAN 2002]. Além de unigramas, essas palavras podem ser agrupadas formando bigramas (duas palavras) ou n-gramas (duas ou mais palavras). Unigramas e bigramas são as principais formas de representação de *tokens* e possuem bons resultados na análise de sentimentos [LIU 2012], tanto na classificação binária [PANG; LEE; VAITHYANATHAN 2002] como multiclasse [PANG e LEE 2005].

Seja a frase “*This cell phone is amazing*”. Na Tabela 2, é exibido um exemplo da representação desta frase em unigramas e bigramas, sem que haja a retirada de nenhuma das palavras em etapas anteriores. Cada n-grama está separado por vírgulas na tabela e a união dos mesmos está representada pelo caractere “_”. Nota-se que a ordem das palavras foi mantida em relação à estrutura da frase inicial e nem todos os n-gramas possíveis estão representados.

Tabela 2. Exemplo de bag-of-words com n-gramas

Unigrama	This, cell, phone, is, amazing
Bigrama	This_cell, cell_phone, phone_is, is_amazing

1.4.3 Técnicas de Seleção de Características

Após a etapa de normalização textual, a fase de Seleção de Características é fundamental para a escolha dos n-gramas para o treinamento de algoritmos de aprendizado [LIU 2012]. Como demonstrado por [PRUSA; KHOSHGOFTAAR; DITTMAN 2015] na análise de dados recolhidos do *Twitter*^{TM12}, a seleção de características pode melhorar

¹¹ Lista de *stopwords* que será utilizada: <http://www.ranks.nl/stopwords>

¹² <http://www.twitter.com>

significativamente o desempenho da classificação. Esta etapa consiste na escolha de n-gramas que serão utilizadas como atributos de treinamento.

Três métodos de seleção de características foram testadas e analisadas: Information Gain, Gain Ratio e Chi-quadrado e estão descritas nas subseções posteriores. Trabalhos como [TANG; TAN; CHENG 2009], [SHARMA e DEY 2012] e [PRUSA; KHOSHGOFTAAR; DITTMAN 2015] fazem uso de alguma técnica de extração de característica.

1.4.3.1 Ganho de Informação

O ganho de informação é uma redução esperada na entropia causada pela divisão dos exemplos de acordo com um atributo qualquer x , na qual entropia é definida como o valor esperado de uma informação [HARRINGTON 2012], considerando-se z o número de classes possíveis que uma informação pode assumir dada pela Equação 1.1:

$$H = \sum_{i=1}^z p(x_i) \log_2 p(x_i) \quad (1.1).$$

Ele mede o número de bits obtidos por meio da predição de uma classe através da presença ou falta de um termo em um documento. Seja t um n-grama, o ganho de informação de um termo é calculado como na Equação 1.2:

$$IG(t) = -\sum_{i=1}^z P(c_i) \log P(c_i) + P(t) \sum_{i=1}^z P(c_i|t) \log P(c_i|t) + P(\bar{t}) \sum_{i=1}^z P(c_i|\bar{t}) \log P(c_i|\bar{t}) \quad (1.2),$$

onde $P(c_i)$ denota a probabilidade de uma classe i ocorrer; $P(t)$ é a probabilidade de um n-grama (atributo) t ocorrer; e $P(\bar{t})$ a é a probabilidade de um n-grama t não ocorrer [TAN e ZHANG 2008].

Em análise de sentimentos, dado um conjunto de n-gramas de uma base de dados com opiniões, na qual duas classes (positivo ou negativo) existem, o IG para cada *token* é calculado com base na Equação 2. Para o problema de RIP com 5 classes, i varia de 1 a 5.

De acordo com a metodologia utilizada, apenas alguns n-gramas são utilizados para treinamento. Estes são escolhidos de acordo com a maior variação do ganho de informação, tanto para classes negativas quanto para classes positivas, isto é, palavras que expressam sentimento negativo, por exemplo, tem maior tendência a serem utilizadas em opiniões nas quais o autor não recomendaria um item.

1.4.3.2 Ganho de Médio de Informação

O ganho médio de informação aprimora o resultado do ganho de informação normalizando a contribuição de todas as características na decisão da classificação final para um documento. Na Equação 1.3, os valores de normalização ou *Split Information* são calculados por meio da informação obtida pela divisão de um documento de treinamento P em v partes, na qual v corresponde a um atributo x [SHARMA, A.; DEY 2012]:

$$SplitInfo(t) = -\sum_{j=1}^v \frac{|P_j|}{|P|} \log \frac{|P_j|}{|P|} \quad (1.3).$$

Por fim, a Equação 1.4 define o ganho médio como:

$$\text{Gain Ratio}(t) = \text{Information Gain}(t) / \text{SplitInfo}(t) \quad (1.4).$$

Assim como no IG, essa fórmula tem como objetivo selecionar palavras que possuem algum sentimento, seja ele positivo ou negativo, e os n-gramas com maior ganho médio são utilizados como atributos.

1.4.3.3 Chi-Quadrado

Este modelo consiste em retirar os n-gramas mais comuns ou os que sejam mais próximos de palavras como “bom” ou “ruim” de um texto. A partir disso, vetores podem ser criados com palavras separadas (unigramas), duas palavras (bigramas) ou n-gramas.

Ele representa a associação entre uma característica e a classe correspondente por meio da Equação 1.5:

$$\text{CHI}(t, c_i) = \frac{N \cdot (AD - BE)^2}{(A+E) \cdot (B+D) \cdot (A+B) \cdot (E+D)} \text{ and } \text{CHI}_{\max} = \max_i(\text{CHI}(t, c_i)), \quad (1.5),$$

onde t é um n-grama e c_i a classe. A é o número de vezes que t e c_i ocorrem simultaneamente; B é o número de vezes que t ocorre sem c_i ; E é o número de vezes que c_i ocorre sem t ; D é o número de vezes que nem c_i nem t ocorrem e; N é o total de documentos [TAN e ZHANG 2008].

Para cada classe, a associação entre um atributo e uma classe é calculada, entretanto, apenas o valor máximo $\text{CHI}_{\max}$ é utilizado, selecionando a classe com maior relação. Na análise textual, t é representado por um n-grama e c são as classes positivo ou negativo na classificação binária ou são as classes referentes as estrelas presentes no problema de inferência de ratings.

1.4.4 Vetorização

Com os n-gramas selecionados pelos métodos de extração de características citados na subseção anterior, a próxima etapa consiste em transformar uma frase em um vetor de características, onde os atributos correspondem aos n-gramas selecionados. Estes atributos são configurados de acordo com frequência dos mesmos em relação a uma opinião. Como exemplo, podemos notar o texto a seguir:

Opinião: *Great Hotel, lovely staff, great location. 8 of us stayed here for 2 nights on a hen party, hotel is close to all bars night clubs, shopping, would definitely stay here again. Hotel is clean and security is great, rooms are really nice and comfortable and have great tv, kitchenette is very handy. Rating: 5.*

Words (15): *great, lovely, worst, location, stay, close, shop, terrible, clean, security, nice, comfortable, handy, bad, good.*

Matriz de representação

4	1	0	1	2	1	1	0	1	1	1	1	0	0	5
---	---	---	---	---	---	---	---	---	---	---	---	---	---	---

Nesse caso, cada posição do vetor corresponde exclusivamente a uma palavra e seu valor é dado pela frequência em uma determinada opinião. As *words* utilizadas acima são apenas exemplos, mas em uma aplicação real, essas palavras são selecionadas pelas técnicas de seleção de características citadas na Seção 1.4.3.

Utilizando o exemplo acima, notamos que a palavra *great* está na primeira posição do vetor. Sua frequência é dada pelo número de vezes que a palavra aparece na frase,

neste caso o número 4. A última posição do vetor corresponde à classe inicial (*rating*) da opinião. Toda a base de dados deve ser configurada seguindo este modelo a fim de criar um grande grupo de exemplos para o treinamento dos algoritmos de aprendizado de máquina supervisionados.

Além da frequência, outro valor pode ser utilizado para preencher cada posição do vetor. O modelo TF-IDF configura os vetores com um peso w_t para um termo t de acordo com a Equação 1.6:

$$w_t = f_t \cdot idf_t = f_t \cdot \log \frac{N}{df_t} \quad (1.6),$$

onde f_t é o número de vezes que t ocorre em uma opinião d ; idf_t é a frequência inversa em um documento do termo t ; N é o total de opiniões e df_t é o número de opiniões que contém t [PALTOGLOU e THELWALL 2010]. Além do trabalho de Paltoglou e Thelwall, que testa várias variantes deste modelo, esta fórmula de representação apresenta bons resultados no trabalho de Martineau e Finin [MARTINEAU e FININ 2009].

Opinião: *Great Hotel, lovely staff, great location.8 of us stayed here for 2 nights on a hen party, hotel is close to all bars night clubs, shopping, would definitely stay here again.Hotel is clean and security is great, rooms are really nice and comfortable and have great tv, kitchenette is very handy. Rating: 5.*

Words (15): *great, lovely, worst, location, stay, close, shop, terrible, clean, security, nice, comfortable, handy, bad, good.*

Matriz de representação

0.887	0.2342	0	0.231	0.887	0.121	0.164	0	0.164	0.123	0.2342	0.421	0	0	5
-------	--------	---	-------	-------	-------	-------	---	-------	-------	--------	-------	---	---	---

1.5 Modelos e Algoritmos de Classificação

Com todo o processo de seleção de características finalizado, criando, por fim, arquivos com atributos quantitativos que são mais facilmente compreendidos e executados por algoritmos de aprendizado, a próxima seção apresenta alguns dos principais modelos e algoritmos nativos utilizados em análise de sentimento para resolver problemas multiclasse. Além disso, ela apresenta dois métodos que utilizam uma forma de classificação binária: *one-versus-one* (OvO) e o *one-versus-all* (OvA), métodos conhecidos como multiclasse adaptado. Os trabalhos que utilizam alguns desses algoritmos estão bem descritos em [LUNARDI; VITERBO; BERNARDINI 2015] e na Seção 1.7 desse trabalho.

1.5.1 Naive Bayes

O algoritmo Naive Bayes é uma variação da teoria de decisão Bayesiana. A probabilidade Bayesiana habilita o conhecimento inicial e a lógica a serem aplicados em declarações desconhecidas [HARRINGTON 2012]. Formalmente, pode-se calcular a probabilidade condicional como na Equação 1.7:

$$P(c | d) = \frac{P(c)P(d | c)}{P(d)} \quad (1.7)$$

Uma variação a teoria bayesiana, o modelo multinomial captura a frequência de uma palavra no conjunto de opiniões [MCCALLUM e NIGAM 1998]. Para associar a

um novo exemplo t uma classe c_i , a classe com maior probabilidade $c^* = \operatorname{argmax} P(c_i|t)$ é considerada. Na Equação 1.8 é mostrado como o cálculo das probabilidades para cada classe $c_i \in c$ é realizado.

$$P_{NB}(c_i | t) = P(c_i) (\prod_{j=1}^D P(t_j | c_i)) \quad (1.8),$$

onde t é um termo, i é o número da classe e D é o conjunto de opiniões.

A partir de um conjunto de termos t de uma opinião, representado pelo vetor w , a distribuição das probabilidades é dada pela Equação 1.9:

$$P_{NB}(c_i | w) = P_{NB}(c_i | t_1 \dots t_n) = \frac{P(w|c_i)P(c_i)}{P(w)} \quad (1.9),$$

no qual n é o número de termos em um vetor w .

Um dos trabalhos iniciais de análise de sentimentos [PANG; LEE; VAITHYANATHAN 2002] utiliza, além do SVM e da Entropia Máxima (EntMax), o Naive Bayes já que este demonstrava bons resultados no problema de categorização de textos. Apesar de simples, o Naive Bayes apresentou bons resultados, superando os outros algoritmos quando treinado com unigramas. Para o problema de multiclasse, Long et. al [LONG; ZHANG; ZHUT 2010] utilizam tanto um modelo de regressão quanto o Naive Bayes, sendo estes treinados com características retiradas de opiniões com uma técnica baseada na complexidade Kolmogorov. Assim como em [PANG; LEE; VAITHYANATHAN 2002], o resultado é satisfatório, chegando a atingir cerca de 12,5% de melhoria de desempenho com os classificadores testados em relação aos trabalhos anteriores.

1.5.2 SVM

Dado um conjunto de dados linearmente separável, caso exista uma linha em um plano que possa separar o conjunto de dados, a linha é chamada de hiperplano separador. A ideia é encontrar o hiperplano que esteja o mais próximo possível dos pontos, sendo que esses pontos estejam o mais distante possível do hiperplano a fim de garantir a melhor robustez do classificador. Isso é chamado de margem. Os pontos mais próximos da margem são chamados de vetores de suporte [HARRINGTON 2012], como pode ser visto na Figura 3.

A ideia principal do modelo de máquina de vetores de suporte é encontrar as margens ótimas em relação a um hiperplano separador h . Essa distância é calculada pela fórmula $u = \vec{w} \cdot \vec{x} - b$, no qual $\vec{w}$ é o vetor normal para o hiperplano, $\vec{x}$ é o vetor de entrada e b é uma constante.

Para o caso linear, a margem é definida pela distância do hiperplano em relação ao vizinho mais próximo dos exemplos positivos e negativos. Maximizar esta margem pode ser expresso por um problema de otimização, no qual a maximização $\frac{2}{||\vec{w}||^2}$ é equivalente a minimizar o problema, conforme a Equação 1.10:

$$L(w) = \frac{||\vec{w}||^2}{2} \quad (1.10),$$

sujeito a $y_n(\vec{w} \cdot \vec{x} - b) \geq 1, \forall n$ no qual x_n é o n -ésimo exemplo de treinamento e y_n é a saída correta do SVM para o n -ésimo exemplo de treinamento.

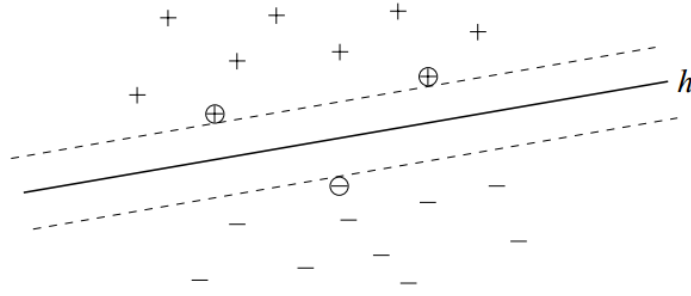


Figura 3. Hiperplano h encontrado, separando dados de treinamento positivos e negativos. Dados circundados são vetores de suporte [JOACHIMS 1998]

Para problemas com a margem suave, variáveis de relaxamento são utilizadas para flexibilizar as restrições do problema de otimização descrito na Equação 1.10. Essas variáveis ξ medem o local de uma amostra em relação as margens. Nesse caso, a Equação 1.10 fica sujeita a $y_n(\vec{w} \cdot \vec{x} - b) \geq 1 - \xi$.

Para problemas não lineares, nem sempre é possível encontrar um hiperplano H para o problema. Para esse tipo de problema, é preciso encontrar uma transformação $\phi(x)$ que não seja linear, de acordo com a Equação 1.11.

$$\phi(x) = \phi_1(x), \dots, \phi_m(x) \quad (1.11),$$

dado m o número de dimensões do problema. Nesse caso, os padrões $\vec{x}$ passam a ser linearmente separáveis e o SVM fica sujeito as restrições $y_n(\vec{w} \cdot \phi(\vec{x}) - b) \geq 1$. Para um conjunto de n padrões $\phi(\vec{x}_n)$, multiplicadores de Lagrange podem ser utilizados. A solução depende apenas do produto $\phi(\vec{x}_i)\phi(\vec{x}_j)$, que pode ser obtido por meio de funções conhecidas como Kernels, como o polinomial mostrado na Equação 1.12.

$$K(\vec{x}_i, \vec{x}_j) = (\delta(\vec{x}_i, \vec{x}_j) + k)^d \quad (1.12).$$

Para problemas multiclasse, são necessários vários classificadores binários que podem ser construídos por meio de técnicas adaptadas, descritas na Seção 1.5.5. Em muitos trabalhos, algumas variantes deste modelo são utilizadas. Isso pode ser notado em [BROOKE 2009] e em [PANG e LEE 2005], no qual o algoritmo Sequential Minimal Optimization (SMO) é o mais indicado para resolver o problema de análise de sentimentos multiclasse já que ele é utilizado para resolver problemas de regressão a partir do SVM. O SMO divide o problema de programação quadrática (PQ) existente no SVM simples, criando soluções menores para o problema de PQ sem utilizar uma matriz de armazenamento extra [PLATT 1998]. Além do SMO, variações do pacote LibSVM¹³ com a função linear sendo utilizada podem ser adaptadas para a classificação multiclasse. Isto se deve ao fato de que, segundo [DUMAIS *et al.* 1998] e [KAESTNER, 2013], o modelo linear é o mais adequado para análise de texto.

Assim como dito em [LIU, 2012], o problema de inferência de *rating* também pode ser considerado um problema de regressão. Isso faz com que variantes do SVM, como o SMO, estejam presentes em trabalhos como [PANG e LEE 2005], [LONG; ZHANG; ZHUT 2010] e [DE ALBORNOZ *et al.* 2011], trabalhos estes que possuem

¹³ <https://www.csie.ntu.edu.tw/~cjlin/libsvm>

bons resultados e são referências na área de inferência de *rating* ou classificação multiclasse.

1.5.3 K-NEAREST NEIGHBORS (KNN)

O k-Nearest Neighbors (kNN) é um método baseado em instâncias que aprende com o simples armazenamento dos dados de treinamento. Quando uma nova instância surge, ele recupera os dados armazenados e classifica essa nova instância [MITCHELL 1997]. A partir dos k vizinhos mais parecidos, ele escolhe o dado com os k mais similares com o que será classificado e atribui uma nova classe a ele [HARRINGTON 2012]. A proximidade dos vizinhos pode ser definida, por exemplo, de acordo com a distância Euclidiana [MITCHELL 1997] demonstrada na Equação 1.13 para dois vizinhos:

$$u = \sqrt{(xA_0 - xB_0)^2 + (xA_1 - xB_1)^2} \quad (1.13).$$

Em Tan e Zhang [TAN e ZHANG 2008], considerando d um documento de teste, a tarefa está em encontrar os k vizinhos entre os outros documentos de treinamento. Na Equação 1.14, a similaridade entre o item d e os outros vizinhos é usada como o peso das classes dos documentos mais próximos, calculado como:

$$\text{score}(d, c_i) = \sum_{d_j \in KNN(d)} \text{sim}(d, d_j) \delta(d_j, c_i) \quad (1.14),$$

no qual $KNN(d)$ representa o conjunto de vizinhos do documento d e c_i uma classe. A função $\text{sim}(d, d_j)$ representa a similaridade entre um documento d o documento de treino d_j . Se d_j pertence a c_i , $\delta(d_j, c_i)$ é igual a 1, senão, é igual a 0. Logo, o documento d deve pertencer à classe que ele possui o maior *score*.

Além de Tan e Zhang, entre os trabalhos relacionados, apenas em Sharma e Dey [SHARMA e DEY 2012] o kNN é utilizado para o problema de classificação de sentimentos binária. Em ambos trabalhos, o kNN apresenta resultado bem inferior quando comparado com os algoritmos SVM e Naive Bayes. Para o RIP, em nenhum dos trabalhos citados nesta pesquisa foi utilizado este algoritmo, o que serviu como motivação para avaliar o desempenho do mesmo neste trabalho. Na ferramenta utilizada neste estudo, o algoritmo kNN é conhecido como IBk (Instance-Based Learning with Parameter k).

1.5.4 Árvores de Decisão

As árvores de decisão são um dos principais métodos de inferência indutiva utilizadas. Elas consistem em um método de aproximação discreta do alvo, na qual a função de aprendizado é representada por uma árvore de decisão, que podem ser representadas como um conjunto de regras *if-then* [MITCHELL 1997].

A tarefa de construir uma árvore de indução consiste em criar uma regra de classificação que pode determinar a classe de objeto a partir dos valores dos seus atributos. Essa regra de classificação é expressa por meio de uma árvore de decisões. As folhas de uma árvore são as classes existentes do problema e os nós internos são os atributos escolhidos no treinamento. A classificação de um novo objeto começa na raiz e para cada atributo uma decisão é tomada a fim de chegar em um novo atributo. Esse processo continua até que a classe apropriada seja encontrada [QUINLAN 1986].

Para o problema de análise de sentimentos multiclasse, os atributos dos nós internos são as características de treinamento selecionados pelos métodos de seleção da Seção 1.4.3 e as folhas representam as classes (para o problema de inferência de *ratings*, cada folha é o valor numérico das estrelas).

O algoritmo ID3, uma implementação de uma árvore de decisões, foi criado para problema nos quais existem muitos atributos e o conjunto de treinamento possui vários exemplos. A ideia básica deste algoritmo é iterativa. Um subconjunto de treinamento é escolhido aleatoriamente e uma árvore é criada a partir dele. Para o restante dos objetos de treinamento são classificados por meio da árvore inicial. Se o restante do conjunto for corretamente classificado, o processo de construção é finalizado. Senão, um conjunto de objetos que não foram corretamente classificados são adicionados ao subconjunto inicial e uma nova árvore é criada. Esse processo pode ser finalizado por meio de um limite de iterações ou até que uma árvore classifique todos os dados de treinamento corretamente. Além do ID3, a variação C4.5 [SHARMA e DEY 2012] é utilizada em análise de sentimentos.

Na pesquisa de Albornoz *et al.* [DE ALBORNOS *et al.* 2011], um modelo de árvore (Functional Tree - FT) é utilizado a fim avaliar o vetor de intensidade de características proposto pelos autores, juntamente com o LibSVM e o algoritmo Logistic. Eles utilizam estes algoritmos com o intuito de prever o *rating* final de uma opinião, atingindo a acurácia de 43,7% para o modelo FT. Entre todos os algoritmos utilizados, a FT obteve o pior desempenho, sendo 3,2% inferior ao modelo Logistic (46,9%).

Chen *et al.* [CHEN *et al.* 2006] criam uma análise visual de opiniões sobre o livro *O Código da Vinci* inspirados em um modelo semelhante a uma árvore de decisões. Além disso, eles utilizam os algoritmos C4.5, SVM e Naive Bayes a fim de selecionar bons termos para a categorização das opiniões utilizadas.

1.5.5 Algoritmos Multiclasse Adaptados

Nessa seção, dois dos principais métodos para resolver problemas multiclasse por meio de divisões binárias são descritos: o *One-vs-One* (OvO) e o *One-vs-All* (OvA). O método OvA cria n divisões, na qual cada etapa do aprendizado é feito comparando um a classe a todas as outras classes. No modelo OvO, cada classe c_i é comparada com outra classe c_k , onde $k, i = 1..n$ e $i \neq k$, dado que n é o número de classes [HSU e LIN 2002].

Em um modelo OvA, a partir da escolha de um classificador (SVM, por exemplo), n classificadores são construídos, isto é, para cada comparação entre uma classe e as demais, um classificador é construído. Para o i -ésimo classificador, os exemplos positivos são todos os pontos da classe i e os exemplos negativos são todos os pontos que não estão na classe i . Seja f_i o i -ésimo classificador, a classificação é dada por meio da Equação 1.15:

$$f(x) = \arg \max_i f_i(x) \quad (1.15).$$

Para o classificador OvO, um modelo classificador também é escolhido, entretanto, cada classe j é comparada com outra classe i . Seja f_{ij} o classificador no qual as classes i são exemplos positivos e as classes j são exemplo negativos. Assumindo que $f_{ij} = -f_{ji}$, a classificação será feita por meio da Equação 1.16:

$$f(x) = \arg \max_i \left(\sum_j f_{ij}(x) \right). \quad (1.16).$$

Em relação ao número de divisões binárias necessárias em cada um desses métodos, no classificador *one-vs-all* ele é dado por i , onde $i=n$, sendo n o número de classes. Já para o algoritmo *one-vs-one*, o número de etapas para a classificação é dado por $\frac{n(n-1)}{2}$, onde n é o número de classes [ALY 2005]. Para o problema multiclasse, os algoritmos SVM citados acima (SMO e LibSVM) utilizam o modelo OvO.

Para um problema com 4 classes {1, 2, 3, 4}, o OvO cria 6 classificadores (1-2, 1-3, 1-4, 2-3, 2-4, 3-4). Para a criação de cada classificador binário, as instâncias de treinamento possuem o rótulo correspondente a cada classificador, isto é, para uma divisão binária 3-4, apenas exemplos de treinamento classificados como 3 ou 4 são utilizados.

Na fase de classificação, a classe escolhida é baseada em uma votação direta dada pelo maior valor de acordo com a Equação 1.16, selecionando a classe com maior número de votos. Exemplificando, dada uma nova instância a , a Tabela 3 e a Tabela 4 mostram uma predição para esse novo dado.

Tabela 3. Votos de cada classificador do modelo OvO

Classificador	$f(a)=$
1-2	2
1-3	1
1-4	1
2-3	2
2-4	2
3-4	3

Tabela 4. Contagem dos votos para cada classe

	Votos para cada classe			
Classe	1	2	3	4
Número de votos	2	3	1	0

Nesse exemplo, a classe 2 é a escolhida como rótulo da nova instância a . Em caso de empate, a escolha é feita aleatoriamente [PIMENTA 2004].

Avaliando o modelo OvA para o mesmo número de classes, 4 classificadores são criados (1 vs {234}, 2 vs {134}, 3 vs {124}, 4 vs {123}). Nesse caso, na fase de treinamento, todas as classes são utilizadas em cada divisão. No processo de classificação, dada uma nova instância a , a predição é dada por meio de Equação 1.15, utilizando uma votação direta distribuída, conforme exibido na Tabela 5. Nesse exemplo, a classe 2 recebe o maior valor (1,999). Dessa forma, a instância a é classificada como 2.

Estas abordagens são comumente utilizadas quando algoritmos SVM são indicados para o problema, por exemplo. Para a avaliação multiclasse, os algoritmos SVM citados acima (SMO e LibSVM) utilizam o modelo de divisões OvO.

Tabela 5. Votos de cada classificador do modelo OvA

$f(a)$	Votos			
Classificador	1	2	3	4

1 vs {234}	1	0	0,333	0,333	0,333
2 vs {134}	Outra	0	1	0	0
3 vs {124}	Outra	0,333	0,333	0	0,333
4 vs {123}	Outra	0,333	0,333	0,333	0
Total		0,666	1,999	0,666	0,666

1.6 Avaliação de Desempenho

Para medir o desempenho dos algoritmos e técnicas citados anteriormente, tipicamente são utilizadas medidas avaliativas, que se baseiam na matriz de confusão. Essas medidas são as mais utilizadas em outros trabalhos como [PANG; LEE; VAITHYANATHAN 2002], [TAN e ZHANG 2008], [GOLDBERG e ZHU 2006] e [GO; BHAYANI; HUANG 2009], seja para a análise multiclasse ou binária.

Uma matriz de confusão para um problema n -classes é uma matriz $n \times n$ [GODBOLE e SARAWAGI 2004] onde o elemento M_{ij} é, para $i=j$, o número de opiniões pertencentes a uma classe i que foram corretamente classificadas e, para $i \neq j$, o número de opiniões de uma classe i que foram erroneamente classificadas em outra classe j . Na Tabela 6 é apresentado um exemplo de uma matriz em que as letras $a-e$ correspondem à escala de *ratings* utilizada (1-5 estrelas), respectivamente.

1.6.1 Acurácia, Precisão e Recall

Com base na matriz de confusão apresentada na Tabela 6, as medidas descritas a seguir são muito utilizadas a fim de medir o desempenho da precisão dos algoritmos, principalmente a acurácia do modelo que mede quão T se aproxima de S , onde T é o conjunto inicial e S é o conjunto com as predições criadas para uma base de dados.

Tabela 6. Exemplo de uma matriz de confusão para o problema 5-classes

a	b	c	d	e	Total	← classificado como
1140	276	61	10	13	1500	a=1
497	502	380	96	25	1500	b=2
132	260	773	281	54	1500	c=3
47	97	228	648	480	1500	d=4
19	36	45	267	1133	1500	e=5
1835	1171	1487	1302	1705	7500	

A acurácia é calculada como:

$$A = \frac{\text{número de exemplos classificados corretamente}}{\text{total de exemplos}}$$

Analisando a Tabela 6, a acurácia final é dada pelo número de exemplos corretamente classificados (1140+502+773+648+1133) dividido pelo número total de exemplos (7500). Desta forma, a acurácia é dada por 0,5594.

A precisão é dada por:

$$P = \frac{\text{número de corretas predições positivas}}{\text{número de predições positivas}}$$

Analisando a Tabela 6, a precisão para a classe a é dada pelo número de corretas predições positivas (1140) dividido pelo número de predições positivas (1835), isto é, o

número de objetos classificados como *a* e que inicialmente eram rotulados como *a* dividido pelo número de exemplos classificados como *a*, sejam eles inicialmente *a* ou não. Desta forma, a precisão é dada pelo valor 0,6212.

O *recall* é dado pela seguinte fórmula:

$$R = \frac{\text{número de corretas predições positivas}}{\text{número de exemplos positivos}}$$

Analisando a Tabela 6, o *recall* para a classe *a* é dada pelo número de corretas predições positivas (1140) dividido pelo número de exemplos positivos (1500), isto é, o número de objetos classificados como *a* e que inicialmente eram rotulados como *a* dividido pelo número de exemplos inicialmente rotulados como *a*. Desta forma, a precisão é dada pelo valor 0,76.

1.6.2 Acurácia Aproximada

O cálculo da acurácia aproximada é definido por Brooke [BROOKE 2009], e esta medida considera aceitável quando uma opinião é classificada com a classe exata ou com a(s) classe(s) vizinhas à classe exata, considerando a escala de *ratings* (1 a 5). A Tabela 7 estende a Tabela 6 para incluir os valores da acurácia aproximada para cada classe. Analisando a classe *b*, tanto opiniões classificadas como *a* ou *c* são aceitáveis e as opiniões inicialmente rotuladas como *b* e classificadas como *d* e *e* são consideradas como erro. Desta forma, se notarmos a acurácia exata da classe *b* (0,335) e a acurácia próxima, concluímos que muitas opiniões da classe *b* foram classificadas como *a* (497) ou *c* (380). Logicamente, o valor da acurácia aproximada é sempre mais elevado do que a acurácia exata.

Tabela 7. Matriz de confusão com acurácia exata e próxima

a	b	c	d	e	Acurácia		← classificado como
					Exata	Aproximada	
1140	276	61	10	13	0,76	0,944	a=1
497	502	380	96	25	0,335	0,919	b=2
132	260	773	281	54	0,515	0,876	c=3
47	97	228	648	480	0,432	0,904	d=4
19	36	45	267	1133	0,755	0,933	e=5
1835	1171	1487	1302	1705	0,559	0,915	-

Essa medida avaliativa também foi utilizada em [PALTOGLOU e THELWALL 2013], no qual os autores consideram não só a acurácia e o erro quadrático médio, mas também exibem o valor do erro absoluto médio e a acurácia aproximada, onde a distância máxima analisada é de uma classe para a classe correta.

1.7 Pesquisas em Análise de Sentimentos

Como foi discutido na seção anterior, os principais métodos de análise de sentimentos podem ser divididos em quatro grandes áreas, de acordo com um modelo semelhante ao de [CAMBRIA *et al.*, 2013]:

- afinidade léxica;
- aprendizado de máquina;
- orientação semântica, e;

- conceitos ou ontologias.

O aprendizado de máquina ou métodos estatísticos consistem na utilização de algoritmos como Naive Bayes e Máquina de Vetores de Suporte a fim de treinar um corpo textual e, a partir do treinamento, classificar novas opiniões. Esses métodos foram anteriormente abordados na Seção 1.5, já que este trabalho tem como foco a proposta de uma técnica que utilize estes algoritmos na análise de sentimentos. Desta forma, esta seção apresenta uma discussão dos principais trabalhos que utilizam algoritmos de classificação para a análise de sentimentos.

Estes trabalhos foram escolhidos com base na importância dos mesmos para a área de análise de sentimentos, levando em consideração os resultados obtidos e as técnicas de extração e algoritmos utilizados. Em alguns casos, algumas destas técnicas e algoritmos não foram citados na seção anterior devido ao grande número de técnicas disponíveis, sendo inviável que todas sejam descritas. Em relação às formas de utilização das opiniões, os principais trabalhos se distribuem em três campos que merecem destaque:

- a classificação em relação à objetividade ou subjetividade;
- a classificação binária, e;
- a classificação multiclasse.

Esses campos serão descritos nas seções abaixo, com destaque para os trabalhos de classificação binária e multiclasse.

1.7.1 Classificação em Texto Objetivo ou Subjetivo

A primeira etapa para realizar a classificação de textos é saber se eles são subjetivos, isto é, contém algum tipo de opinião em relação a uma entidade. Desta forma, tendo uma base de dados que não garanta que existam apenas textos com opiniões subjetivas, uma primeira etapa a ser realizada no processo de análise de sentimentos deve ser separar tais textos em relação à objetividade ou subjetividade. Wiebe e Riloff [WIEBE e RILLOF 2005] desenvolveram um classificador subjetivo usando textos não rotulados para o treinamento. A pesquisa inicia com um processo de busca que utiliza um dicionário de palavras subjetivas para criar os dados de treinamento automaticamente. Esses dados são utilizados para criar um modelo de extração de características e um classificador probabilístico. Finalmente, eles adicionam um mecanismo de autotreinamento que providencia um auxílio aos classificadores, enquanto eles ainda dependem de dados não anotados.

Yu e Hatzivassiloglou [YU e HATZIVASSILOGLOU 2003] utilizaram a similaridade entre sentenças e um classificador Naive Bayes para classificar um texto como subjetivo ou objetivo, baseando-se na afirmativa de que opiniões são mais similares a outras opiniões do que a textos factuais. Eles utilizaram um sistema chamado SIMFINDER para medir a similaridade entre as palavras e frases utilizadas nas diversas sentenças de treinamento. Para realizar a classificação final (objetivo ou subjetivo), os autores utilizaram técnicas de extração como n-gramas, marcadores POS e palavras que possuam algum sentimento. Além disso, a proposta também realizou a classificação binária de uma sentença classificada como subjetiva.

1.7.2 Classificação Binária

Muitos dos trabalhos existentes na área de análise de sentimentos têm como principal objetivo avaliar o desempenho de um ou mais algoritmos de aprendizado, comparando o resultado final, seja por meio da acurácia, tempo ou outras medidas avaliativas. Para isso, são utilizadas bases de dados com avaliações disponíveis na web, com o intuito de avaliar os melhores algoritmos e as melhores técnicas de extração de características. O principal objetivo destes estudos é a classificação em relação à polaridade de uma opinião, isto é, saber se ela é negativa ou positiva; boa ou ruim; recomendada ou não recomendada.

Essa seção discute a grande maioria dos trabalhos referenciados em nossa pesquisa, muitos das quais serviram de base para a metodologia utilizada. Isso se deve ao fato de o problema de análise de sentimentos ser geralmente considerado como um problema de classificação binária [LIU 2012]

Pang *et al.* [PANG; LEE; VAITHYANATHAN 2002] tinham como principal objetivo determinar se uma avaliação é positiva ou negativa utilizando algoritmos de aprendizado. Os autores compararam o desempenho destes algoritmos no problema de mineração de opiniões com o desempenho na classificação feita por humanos e na categorização baseada em tópicos. Os autores mostraram que os algoritmos são melhores na classificação do que humanos, mas seu desempenho não é melhor do que tradicionais métodos de categorização baseado em tópicos (classificação por assunto). Eles utilizaram uma base de dados de avaliações de filmes e pediram para que dois estudantes criassem uma seleção de palavras que indicavam a positividade ou negatividade de uma avaliação. Baseado nessa lista, eles criaram novos vetores de palavras que serão utilizadas pelos algoritmos Naive Bayes, SVM e Entropia Máxima. O desempenho alcançado foi melhor do que as bases formadas por humanos, mas em relação à acurácia de 90% da categorização baseada em tópicos, nenhum dos algoritmos, mesmo quando combinados com bigramas, POS ou a posição de um n-grama no texto conseguiu atingir tal desempenho. O melhor classificador foi o SVM, enquanto a utilização de unigramas mostrou-se mais efetiva em relação às características.

Kang *et al.* [KANG; YOO; HAN 2012] propuseram um novo método para a análise de sentimentos de opiniões sobre restaurantes apresentando duas melhorias no algoritmo Naive Bayes a fim de resolver o problema de balanceamento das acurácias das classificações positivas e negativas. Eles combinaram técnicas de unigramas e bigramas (que incluem tratamento de palavras negativas e utilização de advérbios intensivos) com o algoritmo SVM, o Naive Bayes e as melhorias do Naive Bayes propostas pelos autores. Os autores demonstraram que o Naive Bayes proposto, quando implementado usando bigramas e unigramas, diminui a distância entre a acurácia positiva e a acurácia negativa para 3.6% comparada ao Naive Bayes original e em até 28% em relação ao SVM para opiniões sobre restaurantes.

Xia *et al.* [XIA; ZONG; LI 2011] fizeram um estudo sobre a efetividade do agrupamento de técnicas para tarefas de classificação binária, focando no agrupamento de conjuntos de características e algoritmos de classificação. Eles projetam dois esquemas utilizando POS e dependência sintática e, para cada esquema, utilizam NB, SVM e a Entropia Máxima para a classificação, utilizando a base de dados de filmes disponíveis

em Cornell¹⁴ e o Multi-Domain Sentiment Dataset¹⁵ com avaliações sobre produtos da AmazonTM.

Tan e Zhang [TAN e ZHANG 2008] fizeram um trabalho que apresenta um estudo sobre análise de sentimentos que não usa a língua inglesa, mas sim a chinesa. Os autores utilizam quatro métodos de seleção de características (Informação Mútua, IG, DF e CHI) e cinco algoritmos de aprendizado de máquina (kNN, Naive Bayes, SVM, Winnow e o classificador centroide, estes dois últimos não citados neste estudo) em uma base de dados que contém opiniões sobre três domínios: educação, filmes e eletrodomésticos. Considerando todos os algoritmos de aprendizado, o melhor método de seleção de característica é o Ganho de Informação, que atinge uma média de 88.6% de acurácia. Considerando os métodos de seleção de características, em relação aos algoritmos de aprendizado, o SVM produz a melhor acurácia: 86.8%. Em um dos testes, os autores realizaram o treinamento do SVM em um domínio de eletrodomésticos e utilizaram o conhecimento adquirido para classificar opiniões no domínio de educação. Os autores surpreendentemente obtiveram 0,899 para o valor do MacroF1 para o SVM treinado, ilustrando a possibilidade do uso de modelos treinados em um domínio serem utilizados em outros.

Matsumoto *et al.* [MATSUMOTO *et al.* 2005] analisaram o desempenho do SVM para realizar a classificação binária de avaliações sobre filmes, utilizando dois conjuntos de dados. Os autores extraíram unigramas, bigramas, frequentes subsequências de palavras e sub-árvores dependentes, e usaram tais características para o treinamento de um classificador SVM. Entre os vários testes, eles atingiram 88.3% de acurácia para a primeira base de dados utilizando bigramas, unigramas e árvores de dependência, e 93.7% para o segundo conjunto, utilizando o SVM com bigramas, unigramas, palavras subsequentes e árvores de dependência.

Paltoglou and Thelwall [PALTOGLOU e THELWALL 2010] mostraram que funções de peso adaptadas da Recuperação de Informação (RI) baseadas no cálculo da *tf.idf*[25] e adaptadas para uma configuração particular da análise de sentimentos podem aumentar significativamente o desempenho da classificação. Os autores mostraram que a utilização do SVM adaptado como algoritmo de aprendizado e com essas funções de peso no processo de vetorização, os resultados atingiram até 96% de acurácia. Esse resultado está entre os melhores desempenhos entre os trabalhos relacionados para classificação binária utilizando um algoritmo de aprendizado.

Sharma e Dey [SHARMA e DEY 2012] exploraram cinco métodos de seleção de características em mineração de dados e sete algoritmos de aprendizado de máquina para análise de sentimento em um conjunto de avaliações on-line de filmes. Entre os melhores resultados, o método Gain Ratio (GR), uma variação de IG, foi o que apresentou os melhores resultados. Já em relação aos algoritmos de aprendizado, o SVM possuiu a melhor média de desempenho, considerando as cinco estratégias de seleção, mas o melhor resultado é apresentado pelo Naive Bayes atingindo 90,9% com GR.

Como pode ser observado, muitas das aplicações exploraram novas configurações e novos métodos para melhorar o desempenho dos algoritmos de aprendizado. Xia *et al.* [XIA *et al.* 2011] exploraram métodos agrupados: regras fixas e métodos treinados a fim

¹⁴ Disponível em www.cs.cornell.edu/people/pabo/movie-review-data/.

¹⁵ Disponível em www.cs.jhu.edu/~mdredze/datasets/sentiment/

de melhorar o desempenho dos algoritmos de aprendizado. Sharma and Dey [SHARMA e DEY 2012] fizeram um estudo sobre vários métodos de seleção de características e algoritmos de aprendizado. Paltoglou and Thelwall [PALTOGLOU e THELWALL 2010] utilizaram várias variações do inverso da frequência e atingem acurácia superior a 95%.

Pode ser notado também que alguns trabalhos utilizaram diversos algoritmos de aprendizado, combinados com diversas TEC's, mostrando que em muitos trabalhos houve algum tipo de comparação a fim de obter o melhor algoritmo para o(s) domínio(s) em estudo. Entre as principais TEC's destacaram-se as que analisaram termos e sua frequência em uma opinião. Entre os principais métodos de análise textual estão DF, IG, CHI, unigramas e n-gramas.

Os unigramas e n-gramas foram usados juntamente com outra técnica de extração em algumas pesquisas, com o intuito de selecionar os n-gramas mais importantes e calcular a frequência dos mesmos. Por exemplo, em [PAK e PAROUBEK 2010] foram usados n-gramas para representar palavras que foram obtidas através da análise da frequência de tais palavras chaves, além de marcadores POS. Em [PANG; LEE; VAITHYANATHAN 2002] foram utilizados unigramas, bigramas, POS e adjetivos, considerando em alguns casos a frequência, e em outras a presença de uma palavra. Em [19] foram testados cinco TEC's e sete algoritmos de aprendizado.

Embora exista um grande número de TEC's, esse estudo considera apenas os algoritmos e as TEC's mais utilizados, que foram descritos nas Seções 1.4 e 1.5. Esses métodos geralmente apresentaram bons resultados em outros trabalhos relacionados e aparecem em trabalhos de grande importância na área de análise de sentimentos utilizando algoritmos de aprendizado. Um resumo de todos esses trabalhos está presente na Tabela 8.

1.7.3 Classificação Multiclasse

Os problemas de classificação multiclasse agregam trabalhos que analisam problemas que podem ser divididos em 3 ou mais classes. O problema de inferência de *ratings* é considerado um problema multiclasse, seja em uma escala com 3, 4, 5 ou mais estrelas. Esses problemas também são conhecidos como problemas de escala de multiponto [PANG e LEE 2008].

Em [PANG e LEE 2005], os autores avaliaram a acurácia de humanos em relação à tarefa de determinar o *rating* de um comentário e, posteriormente, eles aplicaram um algoritmo baseado em *metric labing* que, em alguns casos, pode superar o desempenho de algumas versões do SVM e a *baseline* de humanos na classificação de sentimentos em dados com três ou quatro classes.

Goldberg e Zhou [GOLDBERG e ZHU 2006] apresentaram um algoritmo semisupervisionado baseado em grafos a fim de inferir *ratings*, utilizando, em parte, dados não classificados, isto é, não rotulados. Para cada opinião não classificada x , esta foi conectada com outras k vizinhas previamente classificadas. Além disso, a opinião x também foi conectada com suas vizinhas k' não rotuladas. Esse grafo criado com tais relações foi utilizado como treinamento para algoritmos de aprendizado, onde a função $f(x)$ foi utilizada para suavizar o grafo. Como experimento, eles usaram cinco algoritmos de aprendizado baseados em regressão e em *metric labeling*, demonstrando o benefício em utilizar opiniões não rotuladas no problema de inferência de *rating*.

Tabela 8. Resumo dos principais trabalhos em análise de sentimento binária

Autores	Domínio	Seleção de Características	Algoritmos	Acurácia (%)
Pang <i>et al.</i> 2002	Filmes	POS, unigramas, bigramas, posição, adjetivos	NB, EntMax e SVM	82.9 (SVM + unigramas)
Mak <i>et al.</i> 2003	Filmes	IG e DF	Decision Tree, kNN e NB	65 (DT + DF)
Matsumoto <i>et al.</i> 2005	Filmes	Unigramas, bigramas, frequentes subsequências de palavras e sub-árvores dependentes	SVM	93.7 (SVM + unigramas + bigramas, frequentes subsequências de palavras)
Tan e Zhang 2008	Educação, filmes e eletrodomésticos	IG, DF and CHI	Classificador centroide, kNN, NB, Winnow e o SVM	*90.6 (SVM + IG) – Medida Macro F1
Go <i>et al.</i> 2009	Tweets	Palavras com sentimento, bigramas and unigramas	NB, EntMax e SVM	83.0 (EntMax com unigramas + bigramas)
Paltoglou e Thelwall 2010	Filmes	Unigramas e DF – variantes do <i>tfidf</i>	SVM	96.9 (SVM + BM25 <i>tf</i> + variante BM25 delta <i>idf</i>) ^b
Kang <i>et al.</i> 2011	Restaurantes	Unigramas and bigramas	NB, SVM e NB adaptado	81.2 (NB adaptado + unigramas + bigramas)
Xia <i>et al.</i> , 2011	Livros, eletrônicos, DVD's e artigos de cozinha	POS and dependência sintática (Word Relation - WR)	NB, SVM e EntMax	Filme – 86.85 (EntMax + POS) Cozinha – 88.65 (NB + WR)
Sharma e Dey 2012	Filmes	IG, GR, MI, CHI e Belief	NB, SVM, EntMax, DT, kNN, Adaboost e Winnow	90.9 (NB + GR)
Ortigosa <i>et al.</i> 2014	Posts no Facebook	Classificação léxica	J48, NB e SVM	83.27 (SVM + classificador léxico)

Analisando dados do *Twitter*, Pak e Paroubek [PAK e PAROUBEK 2010] coletaram microtextos e os separaram em três classes: sentimento positivo, sentimento negativo e textos objetivos. Esses *tweets* foram selecionados a partir de *emoicons* que apresentassem uma relação com os sentimentos “felizes” ou “tristes”. As TEC’s utilizadas para o treinamento foram n-gramas e a frequência dos mesmos nos *tweets* selecionados. Entretanto, para o treinamento do classificador utilizado (Naive Bayes), eles utilizaram, além de n-gramas, marcadores POS. Como resultado final, eles demonstram que o melhor resultado foi utilizando bigramas, com acurácia chegando a 85%.

Qu *et al.* [QU; IFRIM; WEIKUM 2010] introduziram um novo tipo de *bag-of-opinions*. Seja uma opinião composta de várias frases, cada frase foi assinalada com um *score* e o *rating* foi inferido agregando os resultados dos *scores*. Para determinar o *score*, um método de regressão foi utilizado, no qual o modelo foi inferido baseando-se nos valores de todas as frases por meio de um modelo de n-gramas proposto. Este modelo avalia o *score* de cada unigrama e, por fim, gera um *score* final para uma frase, no qual um unigrama é o foco da frase (raiz), seguido de n-gramas modificadores e negadores. Os autores mostraram que esta técnica supera todos os trabalhos anteriores em uma margem significativa.

Long *et al.* [LONG; ZHANG; ZHUT 2010] propuseram uma nova pesquisa em seleção de opiniões a fim de estimar os *ratings* para serviços em sites utilizando a distância de informação das opiniões por meio da complexidade Kolmogorov. O modelo Kolmogorov associa um valor numérico a cada *string* binária e induz um conceito de similaridade entre tais *strings*. Neste trabalho, a inferência do *rating* foi feita em relação a um atributo do serviço. Isto é, seja um item *A* com vários atributos $a_1, a_2, \dots, a_n$. Para inferir o *rating* para *A*, os autores utilizaram uma combinação dos valores inferidos para cada atributo *a* por meio de classificadores de redes Bayesianas. Este método produziu bons resultados para o problema de análise de sentimentos multiclasse usando qualquer tipo de opiniões, sejam elas compreensíveis (quando estão relacionadas especificamente sobre os atributos de uma entidade) ou não (quando algum atributo não possui uma

opinião) em relação aos atributos utilizados: preço, serviço, quartos e limpeza. Quando o resultado foi estimado para opiniões compreensíveis, a acurácia, entretanto, não chega a 60%.

Albornoz *et al.* [DE ALBORNOZ *et al.* 2011] analisaram o impacto de diferentes características de um produto e o *rating* final. O objetivo é inferir o *rating* com base no *rating* que cada atributo que um produto recebeu em uma determinada avaliação. Para isso, os autores criaram um vetor com a *intensidade dos atributos*, que foi baseado na polaridade e na força da opinião expressada e em outras opiniões associadas a ela, utilizado para o treinamento dos algoritmos de aprendizado. Em relação aos resultados, o algoritmo Logistic (disponível na ferramenta Weka) apresentou o melhor resultado, atingindo 46,9% de acurácia em relação à 5 classes.

Embora também tenham como principal objetivo a classificação multiclasse, Paltoglou e Thelwall [PALTOGLOU e THELWALL 2013] exploraram outros dois tipos de dimensão afetiva para classificar as opiniões: a valência e a excitação. Eles construíram os vetores de características por meio de *tokens* extraídos considerando as duas dimensões afetivas citadas e utilizam um modelo de regressão e uma variação do algoritmo SVM (OVA) para classificar uma opinião em uma escala de sentimento (escala 1-5).

Na Tabela 9, os trabalhos estão organizados destacando o domínio, as TEC's e os algoritmos empregados e a acurácia final. Como citado por Pang e Lee [PANG e LEE 2008], o problema de multiclasse pode ser resolvido por meio da regressão, já que os *rating* são ordinais. Isso pode justificar a escolha da grande maioria dos autores pela utilização do SVM e outros modelos de regressão.

Tabela 9. Resumo dos principais trabalhos em análise de sentimento multiclasse

Autores	Domínio	Técnicas de Extração de Características	Algoritmos	Acurácia (%)
Pang e Lee 2005	Filmes	Frequência de um termo	SVM One-vs-all, Regression and Metric label	59,4 (SVM + vetor de palavras + regressão)
Goldberg e Zhou 2006	Filmes	Modelo semi-supervisionado baseado em grafos	Regressão (SVM), Metric Labeling e PSP	59,2 (regressão+PSP ou regressão)
Pak e Paroubek	<i>Tweets</i>	Frequência, n-gramas e POS	NB	60-80 (NB + bigramas)
Qu <i>et al.</i> 2010	Produtos da Amazon	Bag of opinions	Regressão	-*mostra apenas o erro quadrático médio
Long <i>et al.</i> 2010	Hotéis	Complexidade Kolmogorov + rede bayesiana	SVM	73,1 – 57,3 (Kolmogorov+SVM baseado em atributos)
Albornoz <i>et al.</i> , 2011	Hotéis	Vetor de intensidade das características	Regressão logística, SVM e Árvore funcional	46,9 (vetor+regressão)
Paltoglou e Thelwall 2013	Notícias	Palavras que expressem valência e excitação	Regressão de Vetor de Suporte, SVM (OvA)	51,8 (Excitação + SVM (OvA))

1.7.4 Aplicações da Análise de Sentimentos

Além das pesquisas voltadas para a comparação entre técnicas de aprendizado e seleção de características, pode-se encontrar trabalhos que, a partir do uso das mesmas, apresentam também uma aplicação final. Nos exemplos abaixo, destaque para trabalhos voltados para as áreas de educação e serviços.

Chen *et al.* [CHEN *et al.* 2006] criam uma análise visual de opiniões positivas e negativas do livro “The Da Vinci Code”. Eles utilizam uma ferramenta visual, o TermWatch, para construir uma rede multicamada de termos baseada em associações sintáticas, semânticas e estatísticas. A fim de avaliar os termos que foram selecionados

anteriormente, eles utilizam um modelo preditivo baseado no SVM. Como característica para o treinamento, um conjunto de opiniões positivas e negativas é utilizado. Neste caso, uma opinião é decomposta em três componentes que refletem a presença de termos positivos, negativos e comuns em ambas as categorias.

Mak *et al.* [MAK *et al.* 2003] criaram um sistema de recomendação web utilizando categorização textual de sinopses de filmes armazenadas no IMDB, selecionados do EachMovie database. Primeiramente, eles adaptaram as opiniões a fim de serem utilizadas nos algoritmos, representando-as em vetores, retirando palavras que não possuem informação útil, utilizando valores para cada palavra restante e ranqueando as características do corpus resultante através de três TECs: IG, DF e Informação Mútua. Com essa primeira etapa finalizada, eles utilizaram três algoritmos para construir um classificador para um usuário do sistema: kNN, Decisions Trees e o Naive Bayes. O desempenho final dos algoritmos foi em torno de 60 a 65%, com as árvores de decisão apresentando o melhor resultado, entretanto, a diferença entre os três é pouco significativa.

Em [GO *et al.* 2009], Go *et al.* utilizaram *emoticons* a fim de treinarem opiniões retiradas do Twitter, utilizando algoritmos de aprendizado. Além dos *emoticons*, palavras-chave presentes no site Twittratr¹⁶ que tenham sentimento positivo ou negativo foram utilizadas no treinamento como unigramas e bigramas. Após testes, eles atingiram cerca de 83% de acurácia com o algoritmo Naive Bayes configurado tanto com unigramas como bigramas. Por fim, os autores também disponibilizaram um site, o *sentiment140*¹⁷, onde é possível saber o sentimento sobre algo em relação aos tweets existentes. O site cria uma lista de tweets positivos, negativos e neutros, além de gráficos que mostram qual sentimento é predominante.

Na área de educação, Ortigosa *et al.* [ORTIGOSA *et al.* 2014] construíram um modelo para avaliar postagens no *Facebook*^{TM18} e, a partir da detecção do sentimento habitual do usuário, verificar mudanças emocionais. A aplicação é chamada de SentBuk. Essa informação foi utilizada em sistemas e-learning a fim de recomendar atividades mais adequadas em relação ao humor do estudante em determinado período. Eles construíram um classificador léxico e, quando um grande número de *posts* foi classificado, eles usaram essas mensagens como entrada de treinamento para o algoritmo de aprendizado de máquina. Para realizar os testes eles utilizaram os algoritmos J48, Naive-Bayes e SVM (radial e sigmoide), onde o melhor resultado foi utilizando o algoritmo SVM (sigmoide) com 83% de acurácia.

1.7.5 Dificuldades da Análise de Sentimentos

A grande maioria dos trabalhos de mineração de opiniões existentes tem como foco a mineração de opiniões no idioma inglês. Embora raros, trabalhos como o de Ortigosa *et al.* [ORTIGOSA; MARTÍN; CARRO 2014] e Tang e Zhang [TAN e ZHANG 2008] utilizam os idiomas espanhol e mandarim, respectivamente. Essa falta de trabalhos em alguns idiomas pode dificultar a análise já que os idiomas têm processos de construção diferentes.

¹⁶ <http://twittratr.com/>

¹⁷ <http://www.sentiment140.com/>

¹⁸ <http://facebook.com>

Em relação ao pré-processamento textual, nota-se uma dificuldade na escolha de palavras para treinamento. Isso porque a forma de escrever pode mudar para cada pessoa. Uma expressão que possa indicar um sentimento muito bom para uma pessoa, pode indicar um sentimento nem sempre bom para outra [LIU 2012]. Da mesma forma, uma palavra pode ser utilizada em qualquer classe, logo, o contexto deve ser analisado para compreender o sentimento da mesma. Um exemplo disso são as ironias, muito presentes em avaliações políticas, por exemplo.

Além disso, outra dificuldade está na análise de opiniões que apresentam poucas palavras ou expressões que indiquem algum sentimento. Esse é um caso estudado principalmente em *tweets*. Isso se deve ao fato de um *tweet* ter o número de caracteres limitado a 140. Em alguns casos, como feito em [GO; BHAYANI; HUANG 2009], a utilização de *emoticons* no treinamento é uma opção a fim de melhorar o desempenho da mineração de opiniões. Outro problema está na existência de *herding effects* [WANG e WANG 2014] para a RIP. Isso se deve ao fato de muitas vezes o *rating* final de um usuário não condizer com o comentário. Isso pode acontecer, por exemplo, pelo fato do *rating* ser baseado na média de notas existente no site e não avaliado em relação à opinião por si só. Muitas vezes, pode-se notar que um comentário possui uma avaliação que poderia ter nota máxima (5 estrelas) mas tem o *rating* 4, por exemplo.

1.8 Conclusão

Embora existam muitas técnicas de análise de sentimentos, como foi notado na discussão dos trabalhos e pesquisas da Seção 1.7, as principais técnicas de extração de características foram apresentadas, bem como alguns dos principais algoritmos de aprendizado utilizados na mineração de opiniões que foram selecionados de acordo com a importância e relevância dos trabalhos relacionados.

Esses tipos de classificação (binária e multiclasse) estão amplamente presentes em sistemas e-commerce e são fundamentais para os usuários desse sistema, capturar e compreender de forma correta o sentimento de outros usuários em relação a um item. Por meio dessas técnicas apresentadas, um bom sistema de análise de opiniões pode ser criado para qualquer tipo de domínio.

O problema de classificação binária, modelos de aprendizado podem ser empregado em sistemas semelhantes ao *sentiment140*. Para problemas multiclasse, esse tipo de análise de sentimentos pode ser usado tanto para criar resumos de sentimentos, processo conhecido como sumarização, ou para sistemas de recomendação assistida de *ratings*, no qual estrelas podem ser inferidas para a opinião dos usuários.

Um estudo de caso baseado em avaliações retiradas do site TripAdvisor^{TM19}, disponíveis em²⁰, que foi utilizado na pesquisa de [WANG *et al.* 2010], foi realizado com as técnicas e algoritmos descritos abaixo e pode ser encontrado no link²¹ abaixo.

Referências

Aly, M. (2005) “Survey on Multiclass Classification Methods Extensible Algorithms. Neural Networks”, N. November, P. 1–9.

¹⁹ <http://tripadvisor.com>

²⁰ <http://times.cs.uiuc.edu/~wang296/Data/>

²¹ <http://www2.ic.uff.br/PosGraduacao/Dissertacoes/722.pdf>

- Baccianella, S., Esuli, A. and Sebastiani, F. (2010) “Sentiwordnet 3.0: an Enhanced Lexical Resource for Sentiment Analysis and Opinion Mining”, Proceedings of the Seventh International Conference on Language Resources and Evaluation (LREC’10), V. 0, N. November, P. 2200–2204.
- Beineke, P., Hastie, T., Manning, C. and Vaithyanathan, S. (2004) “Exploring Sentiment Summarization”, Proceedings Of The AAAI Spring Symposium On Exploring Attitude And Affect In Text Theories And Applications, V. 07, P. 1–4.
- Brooke, J. (2009) “A Semantic Approach to Automated Text Sentiment Analysis”, Simon Fraser University, V. 26, N. 4, P. 118.
- Cambria, E., Schuller, B., Xia, Y. and Havasi, C. (2013) “New Avenues in Opinion Mining and Sentiment Analysis. IEEE Intelligent Systems, N. April, P. 15–21.
- Chen, C., Ibekwe-Sanjuan, F., Sanjuan, E. and Weaver, C. (2006) “Visual Analysis of Conflicting Opinions”, IEEE Symposium on Visual Analytics Science and Technology 2006, Vast 2006 - Proceedings, P. 59–66.
- Das, S. R. and Chen, M. Y. (2001) “Yahoo! For Amazon: Opinion Extraction From Small Talk on the Web. Proceedings of the 8th Asia Pacific Finance Association Annual Conference, V. Xxxiii, N. 2, P. 81–87.
- Dave, K., Lawrence, S. and Pennock, D. (2003) “Mining the Peanut Gallery: Opinion Extraction and Semantic Classification of Product Reviews”, Proceedings of the 12th International Conference on World Wide Web, P. 519–528.
- De Albornoz, J. C., Plaza, L., Gervás, P. and Díaz, A. (2011) “A Joint Model of Feature Mining and Sentiment Analysis for Product Review Rating”, Advances in Information Retrieval, p. 55–66.
- Dumais, S., Platt, J., Heckerman, D. and Sahami, M. (1998) “Inductive Learning Algorithms and Representations for Text Categorization”, Cikm ’98: Proceedings of the Seventh International Conference on Information and Knowledge Management, P. 148–155.
- Go, A., Bhayani, R. and Huang, L. (2009) “Twitter Sentiment Classification Using Distant Supervision”, Processing Cs224n Project Report, Stanford, V. 150, N. 12, P. 1–6.
- Godbole, S. and Sarawagi, S. (2004) “Discriminative Methods for Multi-Labeled Classification”, Advances in Knowledge Discovery and Data, V. Lncs3056, p. 22–30.
- Goldberg, A. B. and Zhu, X. (2012) “Seeing Stars When There aren’t Many Stars: Graph-Based Semi-Supervised Learning for Sentiment Categorization”, Proceedings of the First Workshop on Graph Based Methods for Natural Language Processing.
- Harrington, P. (2012) “Machine Learning In Action”, Manning, 2012.
- Kaestner, C. A. A. (2013) “Support Vector Machines and Kernel Functions for Text Processing”, Revista de Informática Teórica E Aplicada, P. 1–7.
- Kang, H., Yoo, S. J. and Han, D. (2012) “Senti-Lexicon and Improved Naïve Bayes Algorithms for Sentiment Analysis of Restaurant Reviews”, Expert Systems with Applications, V. 39, N. 5, p. 6000–6010.
- Konstan, J. A., Miller, B. N., Maltz, D., Herlocker, J. L., Gordon, L. R. and Riedl, J.

- (1997) “Grouplens: Applying Collaborative Filtering to Usenet News”, *Communications of the Acm*, V. 40, N. 3, P. 73–75.
- Kontopoulos, E., Berberidis, C. and Dergiades, T. (2013) “Ontology-Based Sentiment Analysis of Twitter Posts”, *Expert Systems with Applications*, V. 40, N. 10, p. 4065–4074.
- Likert, R. (1932) “A Technique for the Measurement of Attitudes”, *Archives of Psychology*, V. 22, N. 140, P. 1–55.
- Liu, B. (2012) “Sentiment Analysis and Opinion Mining” Morgan and Claypool Publishers, N. May.
- Long, C., Zhang, J. and Zhut, X. (2010) “A Review Selection Approach for Accurate Feature Rating Estimation”, *Proceedings of the 23rd International Conference on Computational Linguistics: Posters*, N. August, p. 766–774.
- Lunardi, A. C., Viterbo, J. and Bernardini, F. C. (2015) “Um Levantamento do Uso de Algoritmos de Aprendizado Supervisionado em Mineração de Opiniões”, *ENIAC - Natal*, RN.
- Mak, H., Koprinska, I. and Poon, J. (2003) “Intimate: A Web-Based Movie Recommender Using Text Categorization”, *Proceedings IEEE/WIC International Conference on Web Intelligence (WI 2003)*, p. 2–5.
- Martineau, J. and Finin, T. (2009) “Delta *tfidf*: An Improved Feature Space for Sentiment Analysis”, *ICWSM*, May, p. 258–261.
- Matsumoto, S., Takamura, H. and Okumura, M. (2005) “Sentiment Classification Using Word Sub-Sequences and Dependency Sub-Trees”, *Proceedings of the 9th Pacific-Asia Conference on Advances in Knowledge Discovery and Data Mining*, V. 05 the 9, p. 301–311.
- McCallum, A. and Nigam, K. (1998) “A Comparison of Event Models for Naive Bayes Text Classification”, *AAAI/ICML-98 Workshop on Learning for Text Categorization*, p. 41–48.
- Mitchell, T. M. (1997) “Machine Learning”..
- Nasukawa, T. and Yi, J. (2003) “Sentiment Analysis: Capturing Favorability Using Natural Language Processing Definition of Sentiment Expressions”, *2nd International Conference on Knowledge Capture*, p. 70–77.
- Ortigosa, A., Martín, J. M. and Carro, R. M. (2014) “Sentiment Analysis in Facebook and its Application to e-Learning”, *Computers in Human Behavior*, v. 31, p. 527–541.
- Pak, A. and Paroubek, P. (2010) “Twitter as a Corpus for Sentiment Analysis and Opinion Mining”, *LREC*, p. 1320–1326.
- Paltoglou, G. and Thelwall, M. (2012) “A Study of Information Retrieval Weighting Schemes for Sentiment Analysis”, *Proceedings of the 48th Annual Meeting of the Association for Computational Linguistics*, n. July, p. 1386–1395.
- Paltoglou, G. and Thelwall, M. (2013) “Seeing Stars of Valence and Arousal in Blog Posts”, *IEEE Transactions on Affective Computing*, v. 4, n. 1, p. 116–123.
- Pang, B. and Lee, L. (2005) “Seeing Stars: Exploiting Class Relationships for Sentiment

- Categorization with Respect to Rating Scales”, In Proceedings of the 43rd Annual Meeting on Association for Computational Linguistics (p. 115-124). Association For Computational Linguistics. v. 3, v. 1.
- Pang, B. and Lee, L. (2008) “Opinion Mining And Sentiment Analysis”, Foundations and Trends in Information Retrieval, v. 2, n. 1, p. 1–135.
- Pang, B., Lee, L. and Vaithyanathan, S. (2002) “Thumbs Up? Sentiment Classification Using Machine Learning Techniques”, Proceedings of the ACL-02 Conference on Empirical Methods in Natural Language Processing-Volume 10, n. July, p. 79–86.
- Pimenta, E. (2004) “Abordagens para a Decomposição de Problemas Multiclasse: os Códigos de Correção de Erro e Saída”, Dissertação da Universidade do Porto, 2004.
- Platt, J. C. (1998) “Fast Training of Support Vector Machines Using Sequential Minimal Optimization”, Advances in Kernel Methods, p. 185 – 208.
- Prusa, J. D., Khoshgoftaar, T. M. and Dittman, D. J. (2015) “Impact of Feature Selection Techniques for Tweet Sentiment Classification”, The Twenty-Eighth International FLAIRS Conference, p. 299–304.
- Qu, L.; Ifrim, G.; Weikum, G. The Bag-Of-Opinions Method For Review Rating Prediction From Sparse Text Patterns. Coling, N. August, P. 913–921, 2010.
- Quinlan, J. R. (1986) “Induction of Decision Trees”, Machine Learning, v. 1, n. 1, p. 81–106.
- Sharma, A. and Dey, S. (2012) “A Comparative Study of Feature Selection and Machine Learning Techniques for Sentiment Analysis”, RAC’S 2012, p. 1–7.
- Tan, S. and Zhang, J. (2008) “An Empirical Study of Sentiment Analysis for Chinese Documents”, Expert Systems with Applications, v. 34, n. 4, p. 2622–2629.
- Tang, H., Tan, S. and Cheng, X. (2009) “A Survey on Sentiment Detection of Reviews”, Expert Systems with Applications, v. 36, n. 7, p. 10760–10773.
- Turney, P. D. (2002) “Thumbs Up Or Thumbs Down ? Semantic Orientation Applied to Unsupervised Classification of Reviews”, Proceedings of the 40th Annual Meeting on Association for Computational Linguistics.
- Wang, H., Lu, Y. and Zhai, C. (2010) “Latent Aspect Rating Analysis on Review Text Data”, Proceedings of the 16th ACM SigKDD International Conference on Knowledge Discovery and Data Mining - KDD’10, p. 783.
- Wang, T. and Wang, D. (2014) “Why Amazon’s Ratings Might Mislead You: The Story of Herding Effects”, Big Data, v. 2, n. 4, p. 196–204.
- Wiebe, J. M. and Riloff, E. (2005) “Creating Subjective and Objective Sentence Classifiers from Unannotated Texts”, Computational Linguistics and Intelligent Text Processing, v. 3406, p. 486–497.
- Xia, R., Zong, C. and Li, S. (2011) “Ensemble of Feature Sets and Classification Algorithms for Sentiment Classification”, Information Sciences, v. 181, n. 6, p. 1138.

Capítulo

2

Keep Calm and Visualize Your Data

Fernanda C. Ribeiro, Bárbara P. Caetano, Melise M. V. de Paula,
Guilherme X. Ferreira, Rafael S. de Oliveira

Abstract

In the last years, information systems have been marked by excessive data available to the user. Information visualization techniques are commonly used to facilitate analysis and understanding of data. However, the constant evolution of technology creates different opportunities for representation of data, which can hinder the development of useful visualizations that are compatible with the purpose of analysis. The purpose here is to present a practical approach that reconciles the literature in information visualization and the various resources available in the construction of graphical representations of data that facilitates the decision-making process.

Resumo

Nos últimos anos, os Sistemas de Informação têm sido marcados pelo excesso de dados disponíveis aos usuários. Técnicas de Visualização de Informação são comumente empregadas para facilitar a análise e o entendimento dos dados. Contudo, a constante evolução da tecnologia cria diferentes oportunidades de representação dos dados, o que pode dificultar a elaboração de visualizações realmente úteis que sejam compatíveis com o propósito da análise. A proposta aqui é apresentar uma abordagem prática que concilia a literatura em visualização de informação e os diversos recursos disponíveis na construção de representações gráficas de dados que sejam facilitadoras para o processo decisório.

2.1. Introdução

O cenário atual é marcado pelo avanço da tecnologia que amplifica as possibilidades de geração de dados. As visualizações de informação vêm sendo cada vez mais abordadas como um meio de viabilizar o entendimento e o uso desses dados.

Historicamente, a visualização de informação vem sendo usada como ferramenta decisiva e fundamental em diferentes processos. Um dos exemplos pioneiros é denominado na literatura como mapa de Snow. Em 1854, após a Revolução Industrial, Londres era marcada pelo crescimento populacional caótico e desorganizado, não havia tratamento de esgoto e a proliferação de doenças era muito comum. Neste mesmo ano, um surto de cólera matou 616 pessoas. John Snow era um físico que contestava a atribuição do miasma como agente da doença. Para refutar esta hipótese, o físico coletou e tabelou os dados sobre a quantidade de mortes em cada endereço. Depois disto, ele elaborou um mapeamento destes dados incluindo a localização das fontes de água distribuídas no local, como ilustrado na Figura 2.1. Ao analisar este mapeamento, foi possível observar que a maioria das mortes se concentrava próxima a uma determinada fonte de água. A partir desta informação, foi constatado que a fonte de água estava contaminada. O surto foi controlado logo após as autoridades terem interditado a fonte [Cairo 2014].



Figura 2.1. Mapa de Snow

Fonte: Retirado de Cairo (2014)

Outro exemplo histórico e clássico foi o mapa elaborado por Charles Minard na campanha russa de Napoleão. No mapa, ilustrado na Figura 2.2, Minard representa o contingente de soldados que cruzaram o rio Niémen no início da campanha (422.000 soldados – representado pela espessura do traço castanho) e a quantidade de soldados que regressaram (apenas 4.000 - indicado pela espessura do traço preto) [Cairo 2014].

Nos dois exemplos apresentados, é possível observar o diferencial quanto à facilidade de análise, proporcionado pelo mapeamento dos dados em uma representação gráfica. Antes de abordar alguns aspectos teóricos e práticos da visualização de informação, é relevante discorrer puramente sobre o termo ou a ação de visualizar, pois é necessário reconhecer algumas nuances que vêm alterando a interpretação da palavra.

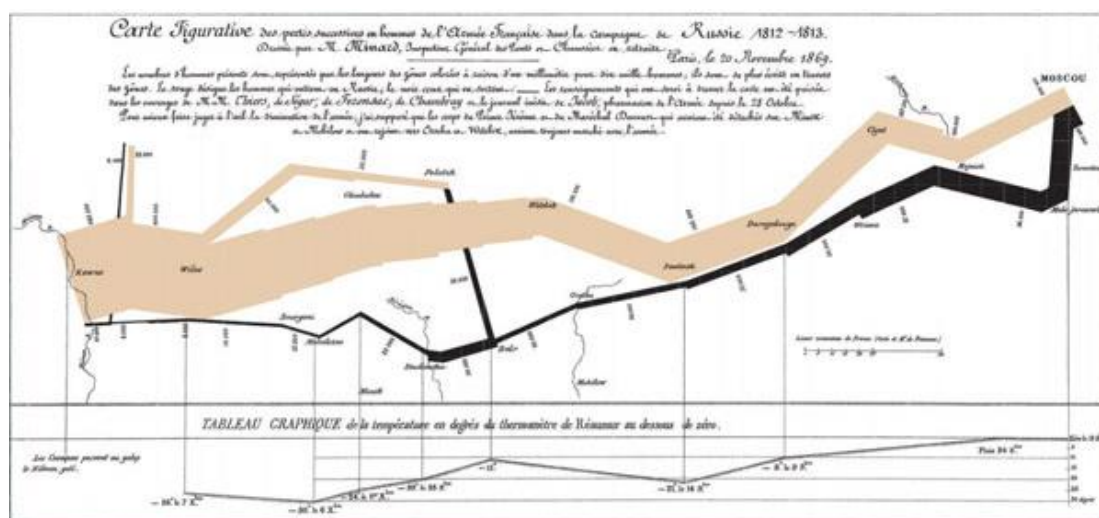


Figura 2.2. Mapa de Minard

Fonte: Retirado de Cairo (2014)

De acordo com Dicio (2016), visualizar significa transformar conceitos abstratos em imagens mentalmente visíveis ou formar uma imagem mental. Nesta definição, estão incluídas quaisquer representações feitas através de imagens podendo ser geradas até mesmo por alucinações. Segundo Ware (2004), visualização pode se referir tanto ao uso de diagramas para transmissão de um sentido quanto ao processamento e transformação de dados em formas gráficas.

No Dicionário Aurélio [Aurélio 2016], encontram-se os seguintes significados para o termo visualização:

- Colocação em evidência de uma maneira material, da ação e dos resultados de um fenómeno.
- Apresentação no ecrã, sob forma gráfica ou alfanumérica, dos resultados de um tratamento de informações.

Comparando os significados da palavra visualização, é possível perceber que o primeiro refere-se ao uso de algum material genérico. Já o segundo significado aborda o uso de uma tela de monitor especificando o recurso a ser utilizado. Essa diferença pode ter sido provocada pelos avanços tecnológicos que fizeram com que o termo deixasse de abordar somente a construção de um modelo mental de forma geral e passasse a descrever um modelo criado através de recursos computacionais [Ware 2004].

Porém, mesmo sendo diferentes do ponto de vista tecnológico, os significados de visualização remetem a um meio para auxiliar o entendimento ou a cognição [Ware 2004], seja colocando em evidência características de um fenômeno ou exibindo resultados de operações com informações. Essa similaridade aponta o principal objetivo da visualização que é aprimorar o sistema cognitivo explorando o sentido da visão e possibilitando uma maior transmissão de informações [Shneiderman 1996] [Ware 2004].

Especificando mais o conceito, a Visualização de Informação (InfoVis) pode ser considerada um tipo de visualização que auxilia no processo de cognição mas está ligada intrinsecamente com a computação. Card and Mackinlay (1997) afirmam, tal qual Shneiderman (1996), que o interesse por esta disciplina de estudo foi ampliado devido ao aumento da disponibilidade de computadores, da capacidade de processamento gráfico e da quantidade de dados.

Na literatura, é possível encontrar várias definições do termo InfoVis como por exemplo:

- InfoVis é uso de representações visuais e interativas dos dados suportadas por computador para ampliar a cognição [Card et al. 1999].
- O objetivo da InfoVis é ampliar o desempenho cognitivo, não apenas para criar imagens interessantes. A InfoVis deve fazer para a mente o que deve fazer os automóveis para os pés [Card 2003].
- A promessa da InfoVis é acelerar nosso entendimento e nossas ações em um mundo de volumes crescentes de informações [Card 2003].

Nestas definições, é possível identificar aspectos importantes do termo, como: o uso do computador, a necessidade de ser esteticamente atraente e interessante e o objetivo principal que é ampliar a capacidade cognitiva do visualizador. Desta forma, a visualização não pode ser considerada estritamente como uma representação gráfica ou o fim de processo, a visualização deve ser entendida como um meio de se chegar a um fim ou uma ferramenta catalizadora do processo de criação do conhecimento.

O objetivo deste minicurso é apresentar conceitos sobre a área de Visualização de Informação além de apresentar um framework que pode ser utilizado desde a concepção da representação do dado até a validação da visualização. Espera-se que ao final do minicurso, o participante tenha adquirido conhecimento não somente das tecnologias abordadas, mas também dos conceitos que fundamentam o processo de construção de uma visualização como uma ferramenta efetivamente útil.

Este capítulo é composto de quatro seções principais. Após da introdução, a Seção 2.2 apresenta o conceito de Visualização de Informação e os aspectos considerados importantes para elaborar uma visualização, como: o Mantra da Visualização, a Teoria de Gestalt e as tarefas perceptuais. Na Seção 2.3 são apresentados alguns exemplos do uso inadequado de algumas técnicas de visualização e também é descrito um framework para elaboração de visualizações. A Seção 2.3 ainda apresenta uma classificação dos recursos e iniciativas relacionados à visualização. Por fim, na Seção 2.4 estão as bibliografias relevantes na área de Visualização de Informação.

2.2. Visualização de Informação

O processo de criação de InfoVis é interdisciplinar por natureza pois requer conhecimento em diversas áreas. Contudo, esta interdisciplinaridade não pode ser considerada apenas uma justaposição de conceitos, as diferentes fontes de conhecimento devem estabelecer uma relação simbiótica e harmônica de forma a promover o benefício

desejado. A Figura 2.3 representa este entendimento. A ilustração pode ser dividida em duas partes e analisada também sob essa perspectiva:

- (a) Os recursos (disciplinas de estudo, ferramentas e conhecimento - tácito ou explícito) necessários para a adequada elaboração de uma visualização de informação.
- (b) Áreas de aplicação nas quais a visualização é parte essencial ou crítica para a obtenção de resultados ou a própria existência da área em questão.

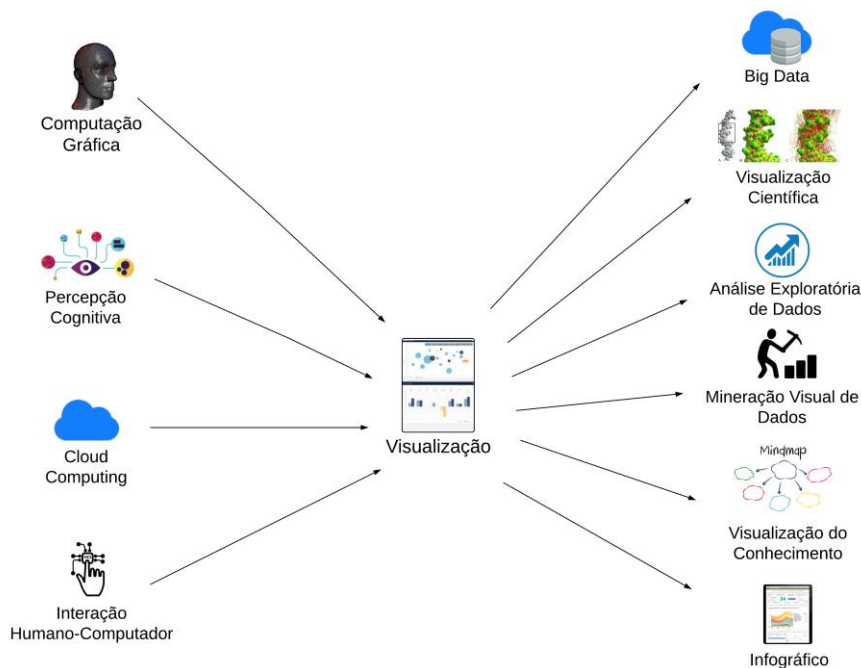


Figura 2.3. Recursos e Áreas de Aplicação da InfoVis

Vale destacar que a Figura 2.3 não pode ser considerada uma representação que ilustra, em sua totalidade, a interdisciplinaridade da InfoVis. Dependendo da abordagem considerada, podem existir outros conceitos/ferramentas que poderiam ser considerados.

Quanto aos recursos, em Freitas et al. (2001), os autores apresentam a InfoVis como sendo uma área de aplicação da computação gráfica cujo objetivo é facilitar a análise de dados. Em Tanahashi et al. (2010), os autores afirmam que, com a propagação da computação em nuvem, a InfoVis torna-se uma aplicação natural desta área pois pode estender o conceito de recuperação da informação à medida que facilita a descoberta de conhecimento. Em Da Silva (2014), o autor afirma que a InfoVis por ser baseada na interação dos usuários com representações visuais e, consequentemente, recebe diversas influências de Interfaces Humano-Computador.

Embora negligenciada, a associação entre a InfoVis e as teorias da Percepção é extremamente relevante para o processo de construção da visualização. Para analisar este aspecto, deve-se fazer três considerações [Alexandre and Tavares 2007]:

- i. A InfoVis é uma ferramenta cognitiva que proporciona a criação do conhecimento através da interpretação da abstração gerada pela percepção visual.

- ii. A InfoVis estimula um processo interpretativo já que a partir de um conjunto de dados, será gerada uma interpretação visual que deverá ampliar a capacidade cognitiva através do estímulo visual.
- iii. A visualização explora principalmente o sentido humano que possui maior aptidão para captação de informação temporal: a visão.

Contudo, a percepção visual e a criação do conhecimento a partir de uma representação gráfica não são fenômenos arbitrários, se baseiam em diversos princípios. Logo, ao elaborar uma visualização de informação é importante considerar estes aspectos.

Em 1996, Shneiderman (1996), ao estabelecer um guia para a construção de visualizações, definiu um princípio básico que foi reconhecido como o Mantra da Visualização: primeiro, visão geral, zoom e filtro, depois, detalhe o que for de interesse. Em 2006, Few (2006) explicou o mantra através de um exemplo. A Figura 2.4 ilustra a variação dos valores de 1.430 ações durante 52 semanas e foi elaborada usando o TimeSearcher 2. Obviamente, a figura não é adequada para comparar valores de ações específicas, mas permite obter várias conclusões a respeito dos dados, por exemplo: a maior concentração de valores está no intervalo entre 7 e 80 dólares e, próximo a 10ª semana, houve um pico de atividade.

Depois de ter uma visão geral dos dados, o próximo passo é focar em algum ponto e examiná-lo com mais detalhes. O foco deve variar em função do interesse da análise. No exemplo dado pelo autor, o foco foi dado no período próximo a 10ª semana. Desta forma, o foco foi dado entre a 8ª e 12ª semana (Figura 2.5). Veja que a barra logo abaixo da figura mostra o foco em relação ao período original de maneira que o visualizador reconheça este período como sendo parte de um todo.

Após o foco, foram selecionadas somente as ações com maiores valores na décima semana. O filtro foi aplicado várias vezes até que se chegasse as 4 ações com maiores valores na 10ª semana (Figura 2.6).

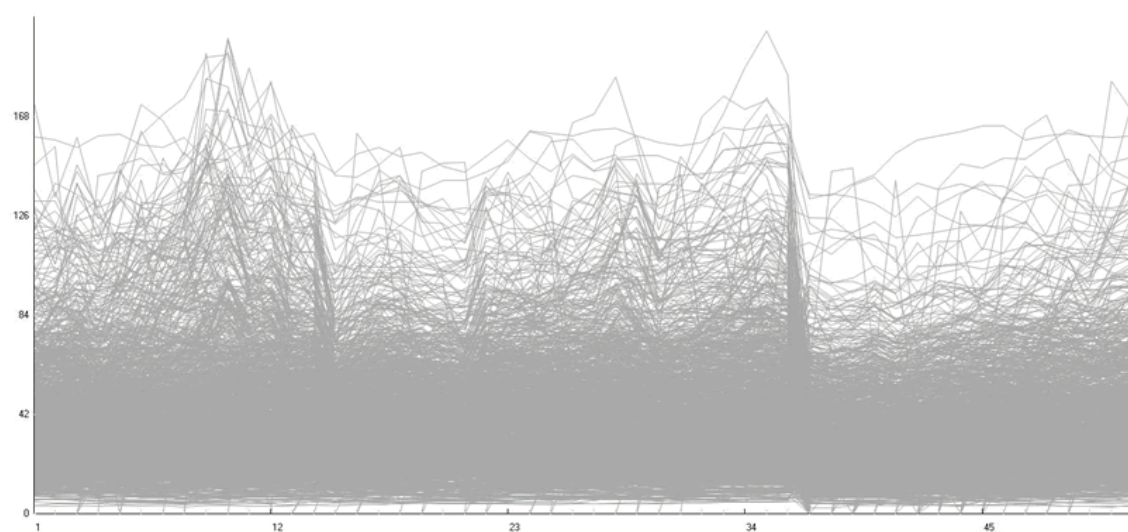


Figura 2.4. Exemplo de InfoVis – Variação de 1430 ações durante 52 semanas

Fonte: Retirado de Few (2006)

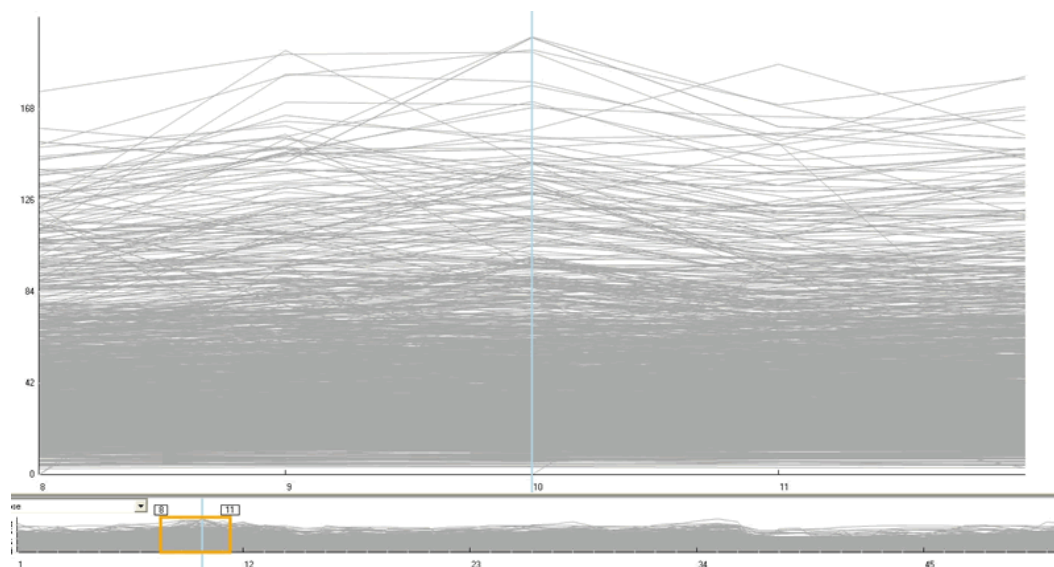


Figura 2.5. Exemplo de InfoVis com Foco – Variação de 1430 ações durante 4 semanas

Fonte: Retirado de Few (2006)

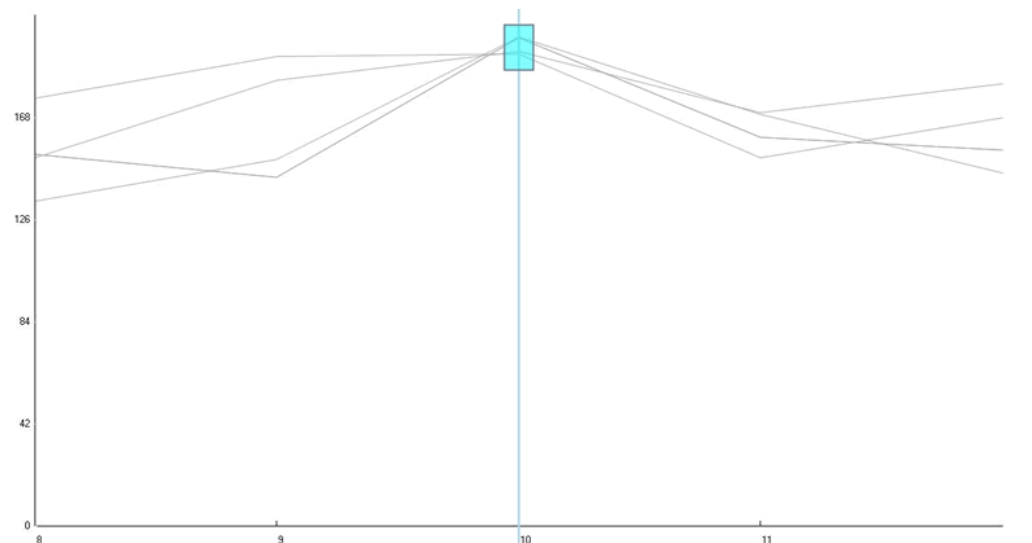


Figura 2.6. Exemplo de InfoVis – Seleção das 4 maiores ações durante uma semana

Fonte: Retirado de Few (2006)

Depois do filtro, são exibidos os detalhes de uma ação em específico (valores por semana e nomes). Na janela de detalhes, são exibidas as ações e seus respectivos valores na semana selecionada (10^a). Na parte inferior da janela de detalhes, são listadas as ações exibidas na janela principal. A Figura 2.7 permite identificar um detalhe interessante dos dados. Veja que na janela de detalhes, são listadas 5 ações, mas aparecem 4 linhas na janela principal. Ao selecionar as demais ações, é possível perceber que uma mesma linha do gráfico representa duas ações diferentes. Esse detalhe pode estimular uma nova investigação iniciando um novo ciclo de análise.

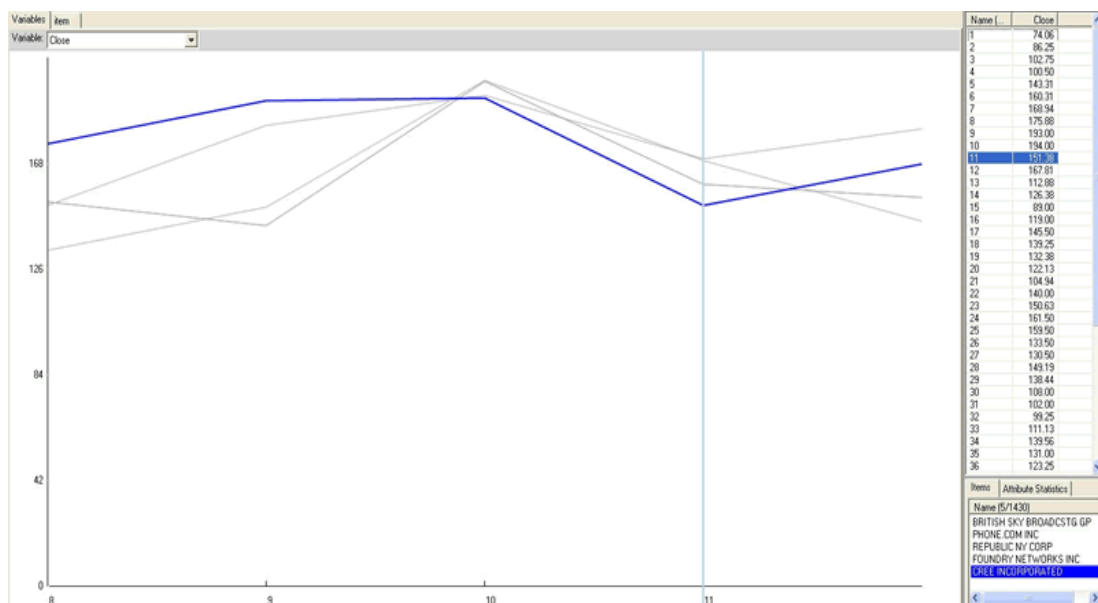


Figura 2.7. Exemplo de InfoVis – Detalhes das ações selecionadas

Fonte: Retirado de Few (2006)

Na Figura 2.8 estão representadas as mesmas cinco ações da Figura 2.7 considerando todo o período de análise (52 semanas) sendo possível identificar que até a 36ª semana, os valores de duas ações foram idênticos. Esta constatação pode sugerir um erro de digitação, por exemplo. Veja que este conhecimento só foi possível obter a partir dos recursos oferecidos pela visualização definidos a partir da sequência lógica estabelecida: o todo, depois as partes. Esta concepção do processo de criação de visualizações está fortemente associada com uma das teorias mais clássicas sobre a percepção: a Teoria de Gestalt.

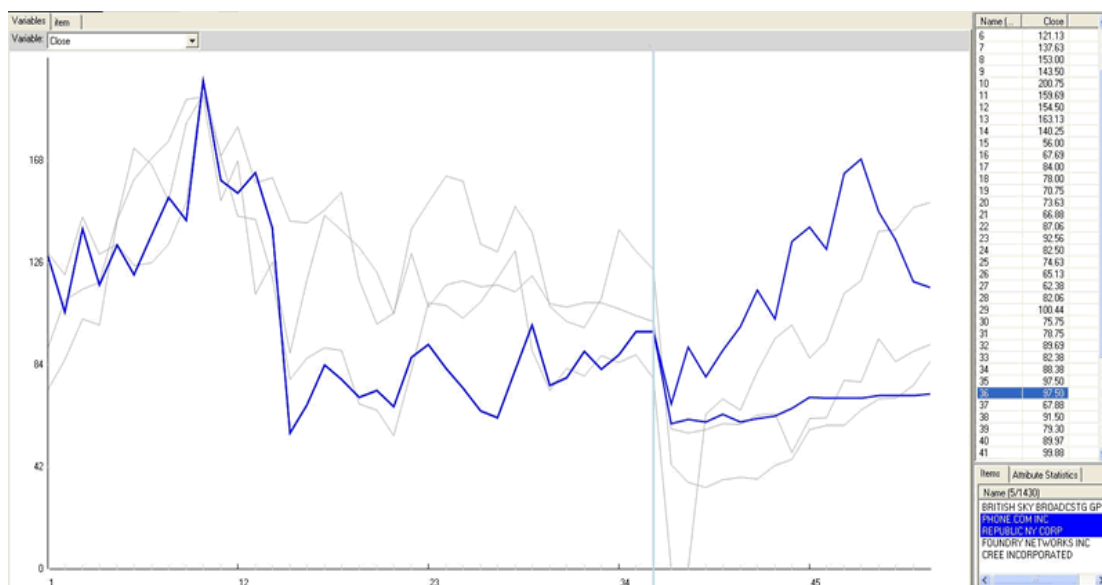


Figura 2.8. Exemplo de InfoVis – Novo ciclo de análise

Fonte: Retirado de Few (2006)

2.2.1. Teoria de Gestalt

O principal axioma da Teoria de Gestalt postula que o todo é mais do que a soma das partes. O conhecimento das partes advém do conhecimento do todo que possui leis próprias que regem suas partes. De acordo com o gestaltismo, a percepção está subordinada a um conceito denominado *Pregnância*. Um objeto é *pregnante* quando o mesmo é fácil de ser percebido através das características estruturais [Alexandre and Tavares 2007]. Segundo Alexandre and Tavares (2007), as características que determinam a *pregnância* de uma imagem são observadas pelas leis da Teoria de Gestalt que são:

- **Proximidade:** elementos que se encontram próximos são percebidos como um grupo [Few 2013], mesmo que não possuam similaridade entre si [Alexandre and Tavares 2007]. Por exemplo, os elementos da imagem à esquerda da Figura 2.9 estão posicionados distantes um do outro e são percebidos como elementos separados. Na imagem à direita, os elementos formam 4 grupos.

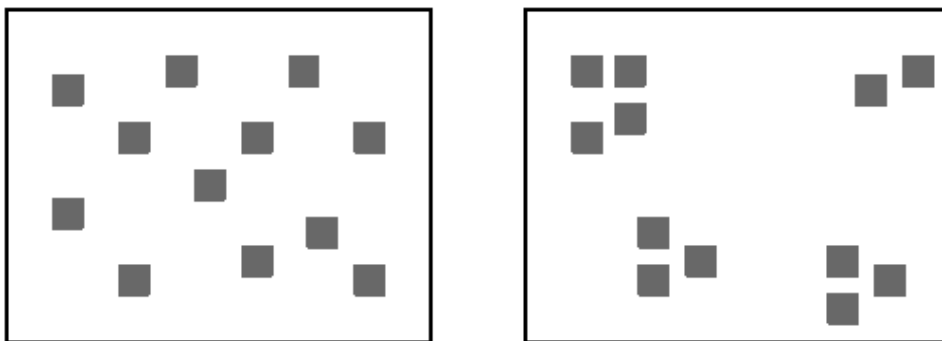


Figura 2.9. Exemplo do princípio de proximidade

Fonte: Retirado de Mazza (2009)

- **Semelhança:** elementos que compartilham características semelhantes (como cor, textura, formato) tendem a ser agrupados [Alexandre and Tavares 2007]. Por exemplo, na Figura 2.10 os objetos de mesmo tamanho (à esquerda) ou com cores diferentes (à direita) são percebidos como pertencentes ao mesmo grupo.

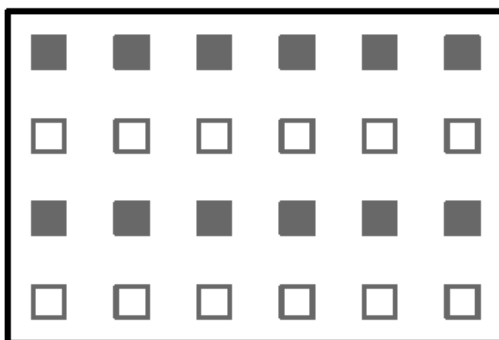


Figura 2.10. Exemplo do princípio de similaridade

Fonte: Retirado de Mazza (2009)

- **Fechamento:** elementos com formas incompletas tendem a ser percebidos como formas completas [Few 2013]. Por exemplo, ao olhar as imagens da Figura 2.11, nós conseguimos enxergar o desenho completo mesmo quando há partes faltando.

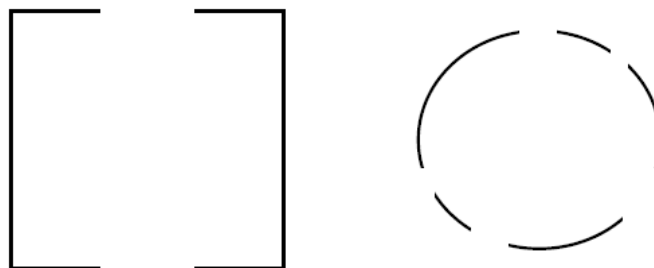


Figura 2.11. Exemplo do princípio de fechamento

Fonte: Retirado de Mazza (2009)

- **Continuidade:** o ser humano tende a perceber uma orientação entre os elementos que parecem construir um fluxo ou um padrão na mesma direção [Alexandre and Tavares 2007]. Segundo Mazza (2009), ao observar a imagem à esquerda na Figura 2.12, nós vemos dois triângulos interrompidos por uma linha horizontal. Não vemos (embora pudesse ser outra possível interpretação da imagem) dois pequenos triângulos, definidos acima da linha, e dois trapézios abaixo [Mazza 2009].

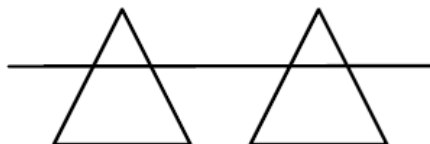


Figura 2.12. Exemplo do princípio de continuidade

Fonte: Retirado de Mazza (2009)

- **Figura/Fundo:** toda imagem pode ser vista como combinação de uma figura e um fundo. A figura distingue-se do fundo por características como: tamanho, forma, cor e posição [Alexandre and Tavares 2007]. Por exemplo, na Figura 2.13, a imagem à esquerda é mais facilmente percebida devido ao alto contraste entre a figura e o fundo.

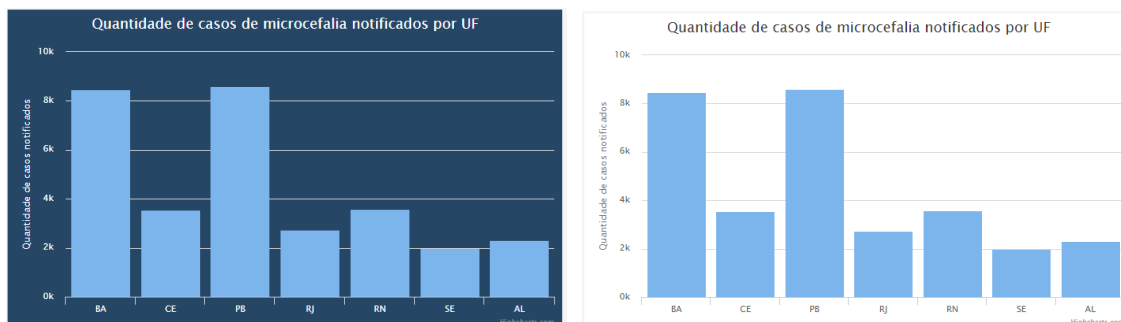


Figura 2.13. Exemplo do princípio de Figura/Fundo

2.2.2. Tarefas Perceptuais

Cleveland and McGill (1984) criaram uma lista de 10 tarefas perceptuais relacionadas à capacidade humana de realizar julgamentos e fazer comparações. De acordo com Cairo (2012), essa lista serve como um guia prático para a escolha da melhor forma gráfica durante o processo de codificação das informações. Essas tarefas estão dispostas em uma escala que varia entre as formas que possibilitam julgamentos mais precisos até julgamentos mais genéricos, conforme mostra a Figura 2.14.

As formas gráficas que possibilitam comparações e julgamentos mais precisos são:

1. Posição ao longo do tempo de uma escala comum
2. Posição ao longo do tempo de escalas não alinhadas
3. Comprimento
4. Direção
5. Ângulo

As formas gráficas que possibilitam comparações e julgamentos mais genéricos são:

6. Área
7. Volume
8. Curvatura
9. Escurecimento
10. Saturação

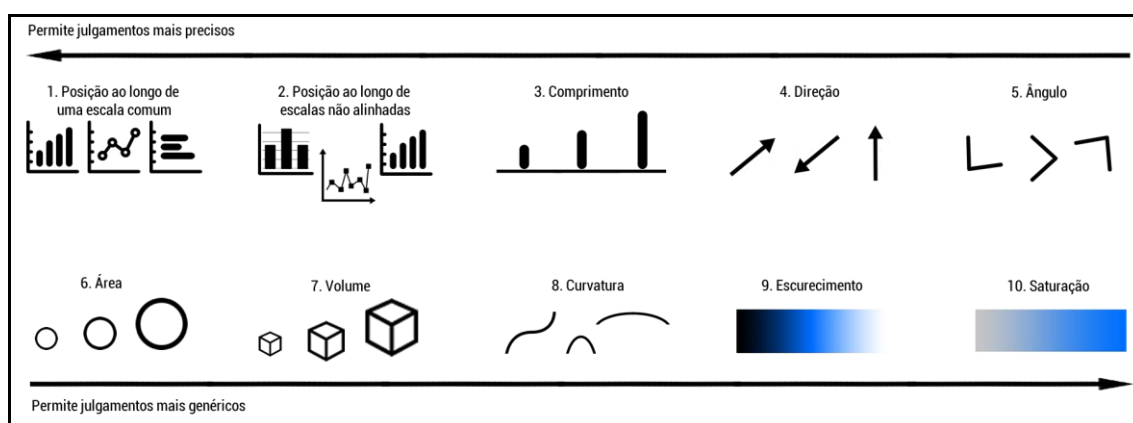


Figura 2.14. Tarefas perceptuais

Fonte: Adaptado de Cleveland and McGill (1984)

De acordo com Cairo (2012), ao utilizar essa escala de Cleveland and McGill (1984), o designer da visualização pode cometer menos equívocos na escolha da forma gráfica. Por exemplo, se o principal objetivo da representação visual for a comparação precisa dos dados, conforme mostra a escala de Cleveland and McGill (1984), utilizar

um gráfico de barras permitiria ao visualizador executar essa tarefa mais facilmente do que utilizando um gráfico de bolhas.

2.3. Visualização na Prática

Atualmente, é notável a diversidade de possibilidades para gerar representações visuais. Segundo Simpson (2014), há uma quantidade crescente de técnicas de visualização criadas para diversas áreas. Essa variedade dificulta a escolha da técnica mais adequada para representar um determinado conjunto de dados. Além disso, é comum encontrar erros não só pela escolha da técnica, mas também pela inadequação da forma como os dados são representados.

Na Seção 2.3.1, serão apresentados alguns exemplos do uso inadequado de algumas técnicas de visualização de maneira que alguns dos erros mais comuns possam ser evitados. Além disso, na Seção 2.3.2, será descrito um framework para elaboração de visualizações de informação baseado na metodologia proposta por Fry (2004).

2.3.1. Exemplos de uso inadequado de técnicas de visualização

Conforme ilustrado na Figura 2.15, o gráfico apresentado pela Globo News sobre o índice de desemprego mostra que, no Brasil, a taxa de desemprego é maior que no Reino Unido, EUA e Alemanha. No entanto, essa informação está errada, pois de acordo com o rótulo da coluna referente ao Brasil, a taxa de desemprego (4,7%) é menor que nesses países que possuem dados acima de 5%. Pode-se supor que a intenção foi destacar os dados referentes ao Brasil. Contudo, o gráfico induz o visualizador achar que a taxa no Brasil é maior, porque a coluna referente ao Brasil é maior que as colunas dos outros países.

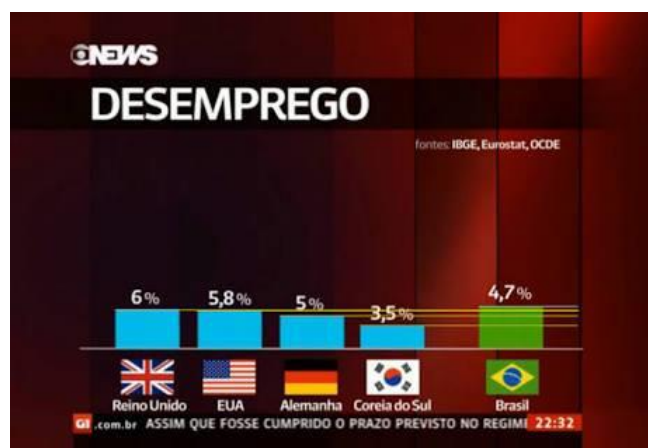


Figura 2.15. Exemplo de uso inadequado - Gráfico de Colunas

Fonte: Retirado de “WTF Visualizations” (2016)

O gráfico da Figura 2.16 mostra dados da anatomia dos convidados do TED Talk, a conferência anual de tecnologia, entretenimento e design. A apresentação dos dados é confusa devido à representação em 3D e em espiral do gráfico de pizza. A largura das fatias que representa as porcentagens é difícil de diferenciar e comparar.

Uma possível solução para melhorar a representação desses dados é remover o 3D, o espiral do gráfico e a textura das cores das fatias, conforme mostra a Figura 2.17.

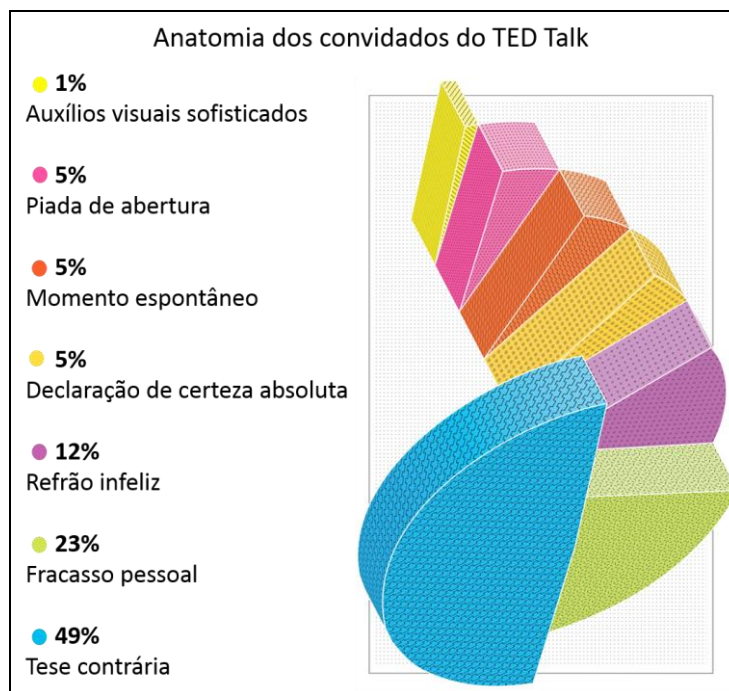


Figura 2.16. Exemplo de uso inadequado – Gráfico de Pizza em Espiral

Fonte: Adaptado de “WTF Visualizations” (2016)

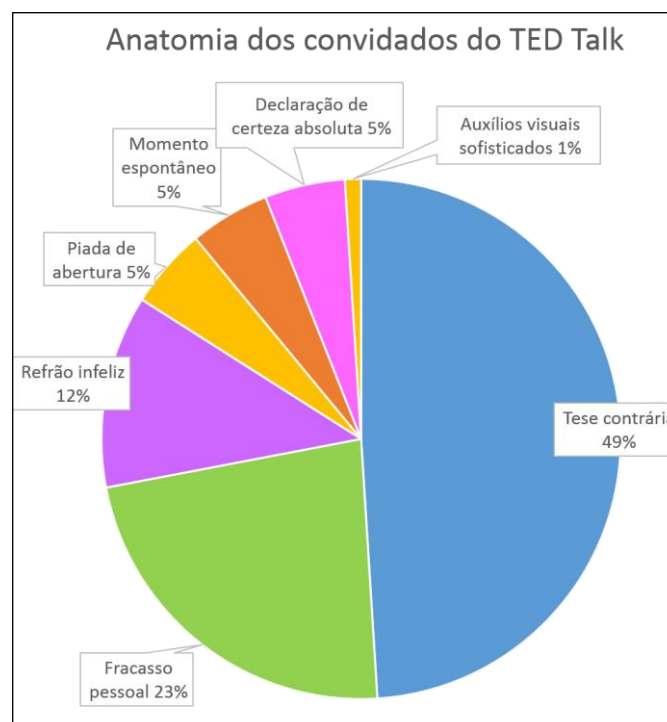


Figura 2.17. Possível solução para o gráfico da Figura 2.16

Os dados apresentados no gráfico da Figura 2.18 mostram as preferências das pessoas em relação à carne de porco. O objetivo foi apresentar qual é a parte mais preferida. O excesso de informação da representação gráfica dificulta e demanda mais tempo para o entendimento dos dados. A textura na borda, a cor do fundo da pizza, as cores no título do gráfico, a legenda redundante, o negrito dos rótulos e o uso do 3D são distrações a mais que não colaboram para o entendimento do dado.

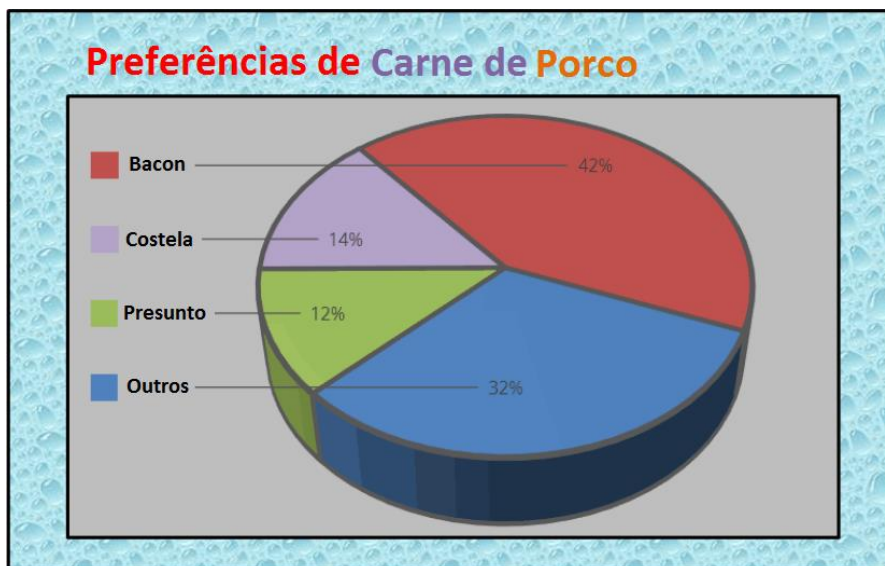


Figura 2.18. Exemplo de uso inadequado - Gráfico de Pizza

Fonte: Adaptado de "Salvaging the Pie" (2016)

A Figura 2.19 apresenta uma possível solução para representar qual é a carne de porco mais preferida. O gráfico de barras neste caso seria mais eficiente por facilitar a comparação. A barra na cor vermelha dá ênfase ao tipo de carne de porco mais preferido. A minimização da quantidade de elementos visuais também auxilia na interpretação rápida do dado.

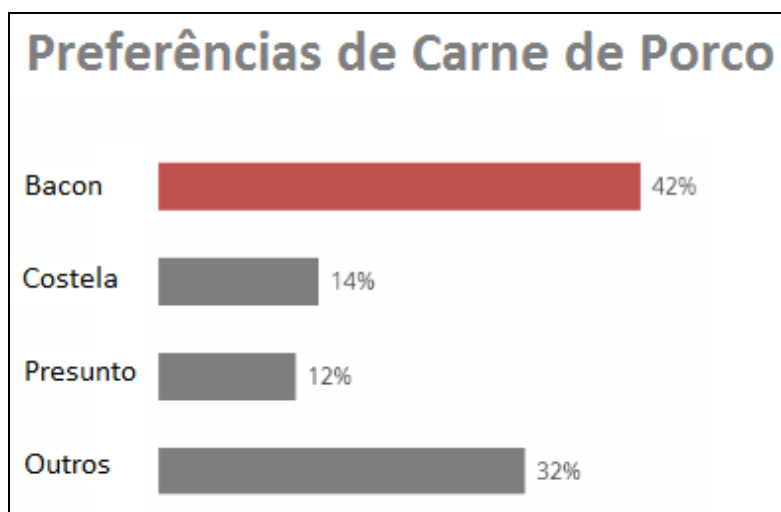


Figura 2.19. Possível solução para o gráfico da Figura 2.18

Fonte: Adaptado de "Salvaging the Pie" (2016)

2.3.2. Framework para criar visualizações

Fry (2004) propôs uma metodologia para auxiliar a criação de visualização de informação que tem como objetivo principal apresentar um método que pudesse atingir um público mais amplo, que não possuísse uma qualificação interdisciplinar, mas que tivesse conhecimento sobre os dados.

Ferreira (2015) identificou a necessidade da adaptação da metodologia proposta por Fry (2004). O trabalho apontou, entre outras melhorias, a inclusão de uma etapa em que fosse possível realizar a análise da qualidade de dados utilizados na visualização. Além disso, diferente da metodologia, a adaptação apresentaria, de forma detalhada, os “passos” que deveriam ser seguidos para atender o objetivo da etapa. A Tabela 2.1 apresenta uma visão geral das etapas.

Tabela 2.1. Visão geral das etapas do framework

Contextualização	Definição do objetivo geral da aplicação e identificação das questões que deverão ser respondidas com os recursos de visualização.
Captura	Obtenção dos dados brutos.
Conversão	Caso os dados capturados não estejam no formato adequado à tecnologia a ser utilizada, é necessário converter estes dados de maneira que possam ser processados.
Limpeza	Tratamento dos dados considerando os aspectos de qualidade de dados que são relevantes para o contexto como, por exemplo, completude dos dados, consistência, duplicatas dentre outros.
Transformação	Gerar novos dados através da transformação dos dados originais, podendo envolver desde a utilização de operações matemáticas simples, como a soma, funções estatísticas, como a média, e técnicas de mineração de dados.
Mapeamento Visual	Consiste na escolha dos elementos gráficos que serão utilizados na aplicação.
Construção Visual	Etapa de desenvolvimento das visualizações.
Interação	Definição e implementação dos recursos de interação do usuário com a visualização.
Validação	Validação da solução elaborada podendo ser utilizados diferentes estratégias.

A esquematização apresentada na Tabela 2.1 deve ser entendida como uma representação simplificada das etapas e não como um modelo rígido. Ainda que a análise deste esquema leve a um entendimento de uma execução sequencial, segundo Fry (2004), as etapas envolvidas na criação de uma visualização de informação tendem a ser executadas de forma cíclica. Além disso, em alguns cenários, algumas etapas podem ser realizadas simultaneamente não sendo possível separar as ações de uma ou outra etapa.

2.3.3. Recursos e iniciativas

A Figura 2.20 apresenta uma classificação dos recursos e iniciativas relacionados à área da visualização de informação. Vale ressaltar que essa classificação não é única e completa.

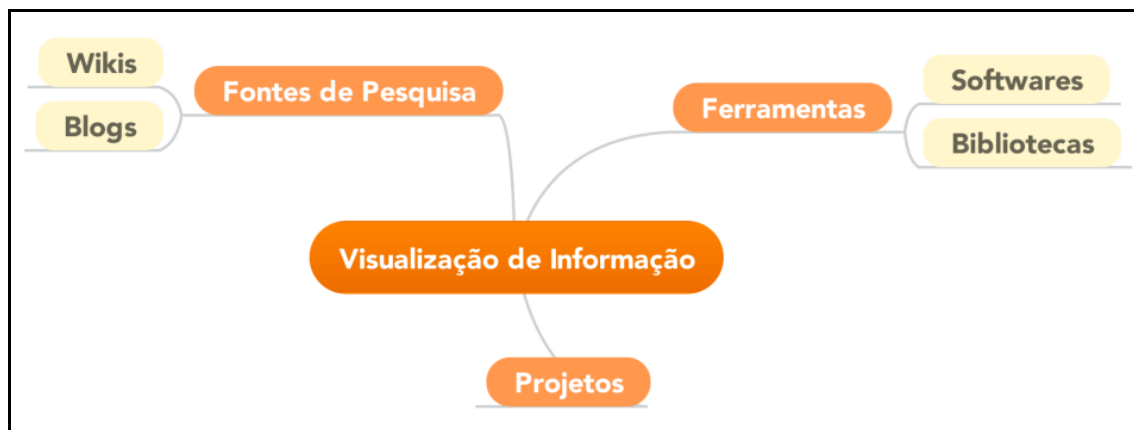


Figura 2.20. Classificação de recursos e iniciativas

Para cada item dessa classificação foi realizado um levantamento dos endereços eletrônicos que são apresentados a seguir.

- **Fontes de pesquisa:**

- **Wikis:**

- Infovis Wiki: <http://www.infovis-wiki.net/index.php>

- Wiki Viz: <http://www.wikiviz.org/>

- **Blogs:**

- The Functional Art: <http://www.thefunctionalart.com/>

- Perceptual Edge: <http://www.perceptualedge.com/>

- Flowing Data: <http://flowingdata.com/>

- Visualising Data: <http://www.visualisingdata.com/>

- WTF Visualizations: <http://viz.wtf/>

- Information is Beautiful: <http://www.informationisbeautiful.net/>

- VizWiz: <http://vizwiz.blogspot.com.br/>

- Eagereyes: <https://eagereyes.org/>

- Visualization Analysis and Design: <http://www.cs.ubc.ca/~tmm/vadbook/>

- Data Visualization: <https://datavisualization.ch/showcases/>

- Information Aesthetics: <http://infosthetics.com/>

- Darkhorse Analytics: <http://www.darkhorseanalytics.com/blog/>

- **Ferramentas:**

- **Softwares:**
 - Tableau: <http://www.tableau.com/>
 - Watson Analytics: <http://www.ibm.com/analytics/watson-analytics/us-en/>
 - Qlik: <http://www.qlik.com/>
- **Bibliotecas:**
 - D3.js: <https://d3js.org/>
 - D3plus: <http://d3plus.org/>
 - Protovis: <http://mbostock.github.io/protovis/>
 - HighCharts: <http://www.highcharts.com/>
 - JavaScript InfoVis Toolkit: <http://philogb.github.io/jit/>
 - Chart.js: <http://www.chartjs.org/>
 - Chartist.js: <https://gionkunz.github.io/chartist-js/>
 - Prefuse: <http://prefuse.org/>
- **Projetos:**
 - The Data Visualization Catalogue: <http://www.datavizcatalogue.com/>
 - DataViva: <http://pt.dataviva.info/>
 - VisPublica: <http://vispublica.gov.br/>
 - Helpmeviz: <http://helpmeviz.com/2014/02/05/have-data-will-visualize/>
 - Seeing Data: <http://seeingdata.org/>
 - Brightpoint: <http://www.brightpointinc.com/>
 - Visual Telling: <http://www.visual-telling.com/>
 - Labvis: <http://labvis.eba.ufrj.br/>

2.4. Bibliografia

Na Tabela 2.2 são apresentadas bibliografias relevantes na área de Visualização de Informação.

Tabela 2.2. Bibliografias na área de Visualização de Informação

Livro: The Visual Display of Quantitative Information
Autor: Edward Tufte
Editores: Graphics Press
Ano: 1983

Livro: Visualize This: The FlowingData Guide to Design, Visualization, and Statistics Autor: Nathan Yau Editora: John Wiley & Sons Ano: 2011
Livro: Information Dashboard Design: Displaying Data for at-a-glance Monitoring Autor: Stephen Few Editora: Analytics Press Ano: 2013
Livro: Data Visualization: a successful design process Autor: Andy Kirk Editora: Packt Ano: 2012
Livro: The Visual Organization: Data Visualization, Big Data, and the Quest for Better Decisions Autor: Phil Simon Editora: Wiley Ano: 2014
Livro: The Functional Art: An introduction to information graphics and visualization Autor: Alberto Cairo Editora: New Riders Ano: 2012
Livro: The Wall Street Journal Guide to Information Graphics: The Dos and Don'ts of Presenting Data, Facts, and Figures Autor: Dona M. Wong Editora: W. W. Norton & Company Ano: 2010
Livro: Visualizing Data: Exploring and Explaining Data with the Processing Environment Autor: Ben Fry Editora: O'Reilly Ano: 2008

Livro: Show Me the Numbers: Designing Tables and Graphs to Enlighten Autor: Stephen Few Editora: Analytics Press Ano: 2014
Livro: Readings in Information Visualization: Using Vision to Think Autor: Stuart K. Card, Jock Mackinlay e Ben Shneiderman Editora: Morgan Kaufmann Ano: 1999
Livro: Beautiful Visualization: Looking at Data through the Eyes of Experts Autor: Julie Steele e Noah Iliinsky Editora: O'Reilly Media Ano: 2010
Livro: Data Flow: Design graphique et visualisation d'information Autor: Robert Klanten Editora: Thames & Hudson Ano: 2009
Livro: Data Flow 2: Visualizing Information in Graphic Design Autor: Robert Klanten, N. Bourquin e S. Ehmann Editora: Die Gestalten Verlag Ano: 2010
Livro: Information Graphics: A Comprehensive Illustrated Reference Autor: Robert L. Harris Editora: Oxford University Press Ano: 2000
Livro: Visual Meetings: How Graphics, Sticky Notes and Idea Mapping Can Transform Group Productivity Autor: David Sibbet Editora: Wiley Ano: 2010

Livro: Information Dashboard Design: The Effective Visual Communication of Data Autor: Stephen Few Editora: O'Reilly Media Ano: 2006
Livro: Information Visualization: Perception for Design Autor: Colin Ware Editora: Morgan Kaufmann Ano: 2004
Livro: Visual Explanations: Images and Quantities, Evidence and Narrative Autor: Edward R. Tufte Editora: Graphics Press USA Ano: 1997
Livro: Envisioning Information Autor: Edward R. Tufte Editora: Graphics Press USA Ano: 1990
Livro: The Elements of Graphing Data Autor: William S. Cleveland Editora: Hobart Press Ano: 1994
Livro: Visualizing Data Autor: William S. Cleveland Editora: Hobart Press Ano: 1993
Livro: Visualization of Time-Oriented Data Autor: Wolfgang Aigner, Silvia Miksch, Heidrun Schumann e Christian Tominski Editora: Springer Ano: 2011

Livro: How to Lie with Statistics Autor: Darrell Huff Editora: W. W. Norton & Company Ano: 1993
Livro: Cartographies of Time: A History of the Timeline Autor: Daniel Rosenberg Editora: Princeton Architectural Press Ano: 2012
Livro: Information is Beautiful Autor: David McCandless Editora: Hardcover Ano: 2012

Referências

- Alexandre, D. S. and Tavares, J. M. R. S. (2007). Factores da percepção visual humana na visualização de dados. Congresso de Métodos Numéricos em Engenharia, XXVIII CILAMCE-Congresso Ibero Latino-Americano sobre Métodos Computacionais em Engenharia,
- Aurélio (2016). Dicionário Português - Dicionário do Aurélio Online. <https://dicionariodoaurelio.com/>, [accessed on Apr 21].
- Cairo, A. (2012). The Functional Art: An introduction to information graphics and visualization. New Riders.
- Cairo, A. (2014). The Functional Art: An Introduction to Information Graphics and Visualization: Explaining Snow, Minard, and Nightingale. . <http://www.thefunctionalart.com/2014/05/explaining-snow-minard-and-nightingale.html>, [accessed on Apr 5].
- Card, S. (2003). Information visualization. In: Jacko, J. A.; Sears, A.[Eds.]. . The Human-computer Interaction Handbook. Hillsdale, NJ, USA: L. Erlbaum Associates Inc. p. 544–582.
- Card, S. K. and Mackinlay, J. (1997). The structure of the information visualization design space. In Proceedings of the 1997 IEEE Symposium on Information Visualization (InfoVis '97). , INFOVIS '97. IEEE Computer Society. <http://dl.acm.org/citation.cfm?id=857188.857632>, [accessed on May 30].
- Card, S. K., Mackinlay, J. D. and Shneiderman, B. [Eds.] (1999). Readings in information visualization: using vision to think. San Francisco, CA, USA: Morgan Kaufmann Publishers Inc.

- Cleveland, W. S. and McGill, R. (1984). Graphical perception: Theory, experimentation, and application to the development of graphical methods. *Journal of the American Statistical Association*, v. 79, n. 387, p. 531–554.
- Da Silva, C. G. (2014). Visualização de Informação: Introdução e Influências de IHC. In *Companion Proceedings of the 13th Brazilian Symposium on Human Factors in Computing Systems*, IHC '14. Sociedade Brasileira de Computação. <http://dl.acm.org/citation.cfm?id=2738165.2738200>, [accessed on Apr 21].
- Dicio (2016). Dicionário Online de Português. <http://www.dicio.com.br/>, [accessed on Apr 21].
- Ferreira, G. X. (2015). Explorando o Processo de Criação de Visualização de Informação no Ambiente de Dados Abertos Governamentais. Universidade Federal de Itajubá.
- Few, S. (2006). The Surest Path to Visual Discovery by Stephen Few - BeyeNETWORK. <http://www.b-eye-network.com/view/2674>, [accessed on Apr 21].
- Few, S. (2013). Data Visualization for Human Perception. *The Encyclopedia of Human-Computer Interaction*, 2nd Ed.,
- Freitas, C. M. D. S., Chubachi, O. M., Luzzardi, P. R. G. and Cava, R. A. (2001). Introdução à visualização de informações. *Revista de informática teórica e aplicada*, v. 8, p. 143–158.
- Fry, B. J. (2004). Computational information design. Massachusetts Institute of Technology.
- Mazza, R. (2009). *Introduction to Information Visualization*. 1. ed. Springer Publishing Company, Incorporated.
- Salvaging the Pie (2016). <http://www.darkhorseanalytics.com/blog/salvaging-the-pie>, [accessed on Apr 23].
- Shneiderman, B. (sep 1996). The eyes have it: a task by data type taxonomy for information visualizations. In , *IEEE Symposium on Visual Languages*, 1996. Proceedings.
- Simpson, (primeiro) (2014). Survey Analysis: A Beginner's Guide. DataHero Official Blog. <https://datahero.com/blog/2014/04/03/survey-analysis-a-beginners-guide/>, [accessed on Jul 4].
- Tanahashi, Y., Chen, C. K., Marchesin, S. and Ma, K. L. (nov 2010). An Interface Design for Future Cloud-Based Visualization Services. In *2010 IEEE Second International Conference on Cloud Computing Technology and Science (CloudCom)*.
- Ware, C. (2004). *Information visualization: perception for design*. 2. ed. San Francisco, CA: Morgan Kaufman.
- WTF Visualizations (2016). <http://viz.wtf/?og=1>, [accessed on Apr 23].

Capítulo

3

Modelos de Negócios Inovadores na Era da Computação em Nuvem

Ricardo Batista Rodrigues, Carlo M. R. Silva, Vinicius C. Garcia, Silvio R. L. Meira

Abstract. *Business models are few understood in our generation. The initial image is a complex and laborious, seen by many as unnecessary. Business models have become an essential tool for entrepreneurs. With technological developments emerged several businesses, some of these businesses are in the cloud or are cloud-based. In this scenario, this work presents the main concepts of cloud computing and business models.*

Resumo. *Modelos de negócios são poucos compreendidos em nossa geração. A imagem inicial é de algo complexo e trabalhoso, visto por muitos como desnecessário. Os modelos de negócios se tornaram uma ferramenta fundamental para empreendedores. Com a evolução tecnológica surgiram diversos negócios, parte desses negócios estão na nuvem ou são baseados em nuvem. Diante desse cenário, este trabalho apresenta os principais conceitos de computação em nuvem e modelos de negócios.*

3.1 Introdução

O desenvolvimento tecnológico e o atual cenário de economia moderna tornaram o mercado exigente e altamente competitivo. No que diz respeito a *software*, cada vez mais o mercado e os clientes esperam por resultados rápidos e concretos, o que exige que os empreendedores sejam ágeis, ousados e eficientes no desenvolvimento de produtos e serviços inovadores [Pressman 2011]. A evolução dos meios de comunicação via *Internet* por diversos tipos de dispositivos, trouxe um novo cenário para o empreendedorismo na área de tecnologia. Cada vez mais surgem *startups* voltadas para produtos e serviços baseados em nuvem, essa tendência deve continuar nos próximos anos [Rodrigues 2015].

O empreendedorismo junto a todo o ecossistema de computação em nuvem ganhou força dentro das universidades, que se consolida como instituição de fomento à inovação em novos negócios. Podemos considerar ainda que o contexto atual é bastante favorável ao crescimento do empreendedorismo e da inovação nas universidades, enquanto alternativa as oportunidades provindas da expansão da computação em nuvem e do acesso a *Internet*. É em busca de novos conhecimentos e soluções inovadoras que os países e universidades vêm investindo em centros de pesquisas e tecnologia, para que a sustentação desse tipo de empreendimento cresça e possa desenvolver soluções que resolvam problemas reais do dia a dia. Ao mesmo

tempo, grandes empresas estão investindo em tecnologias baseadas em computação em nuvem. A cada dia surgem novos negócios baseado em nuvem, muitos não possuem escritórios físicos e são em sua totalidade baseados em nuvem [Rodrigues 2015].

Este trabalho apresenta os principais conceitos sobre computação em nuvem, empreendedorismo e modelos de negócios. Serão detalhadas algumas técnicas e ferramentas que são utilizadas com frequência por empreendedores, membros e *coachs* de novos negócios (*startups*). Esse trabalho tem como objetivo apresentar como essas técnicas e ferramentas podem ser utilizadas de forma simples e prática no desenvolvimento de novas ideias inovadoras baseadas em tecnologias da informação.

Este trabalho está organizado da seguinte forma, na Seção 2 são apresentados os principais conceitos sobre computação em nuvem; na Seção 3 são apresentados os principais conceitos sobre modelos negócios; na Seção 4 são apresentados os principais conceitos sobre inovação e startups; na Seção 5 são apresentados conceitos de como torna uma ideia em negócio; Na Seção 6 são apresentadas tendências de mercado; e na Seção 7 são apresentadas considerações finais.

3.2 Computação em Nuvem

O *National Institute of Standards and Technology* (NIST) define computação em nuvem como um modelo que permite que um conjunto de recursos computacionais possam ser fornecidos sob demanda de forma a permitir que os mesmos sejam fornecidos e liberados rapidamente com o mínimo de esforço de gestão ou interação do fornecedor [Mell and Grance 2009]. [Vaquero et al. 2008] definem computação em nuvem como um grande conjunto de recursos virtualizados (*hardware*, plataformas de desenvolvimento e/ou serviços) facilmente usáveis e acessíveis. A arquitetura da computação nas nuvens é comumente representada em camadas: *Software* como serviço, *Plataforma* como Serviço e *Infraestrutura* como serviço. Nesse modelo cada camada está construída sobre os serviços oferecidos pela camada de baixo [Lenk et al. 2009], conforme apresentado na Figura 1.

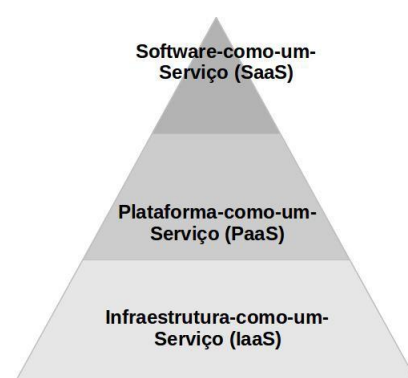


Figura 1. Arquitetura em camadas para computação em nuvem [Machado 2013]

- **IaaS:** é o fornecimento de recursos computacionais como: capacidade de processamento, armazenamento e conectividade. Tomando armazenamento como um exemplo, quando o usuário usa um serviço de armazenamento em computação em nuvem, ele somente paga a parte que foi efetivamente consumida sem comprar nenhum disco rígido ou até mesmo sem nem conhecer a localização dos seus dados armazenados. Um exemplo de IaaS é o *Amazon Web Services* (*Amazon AWS*),

principalmente, através do *Amazon Elastic Compute Cloud (Amazon EC2)*;

- **PaaS:** fornece a infraestrutura e dá suporte a um conjunto de interfaces de programação para aplicações em nuvem. É o meio entre o hardware e as aplicações. Um bom exemplo de PaaS é a plataforma de serviços do Microsoft Azure³;
- **SaaS:** em SaaS o objetivo é substituir as aplicações que rodam nos computadores. Ao invés de comprar o *software* por um preço relativamente alto, o usuário paga pelo que é usado do SaaS, o que pode gerar economia, observando que o usuário só pagará pelo que usar. Um dos principais problemas é a baixa taxa de transmissão de dados na rede em determinadas localidades o que é fatal para algumas aplicações, como sistemas de tempo real. Um exemplo muito conhecido de SaaS é o *Dropbox*⁴;

A computação em nuvem traz três novos aspectos em *Hardware* [Vogels 2008]:

1. A ilusão de recursos computacionais infinitos disponíveis sob demanda, eliminando assim a necessidade dos usuários planejarem muito à frente para provisionamento de recursos;
2. A eliminação de um compromisso antecipado por parte dos usuários da nuvem, permitindo que as empresas comecem pequeno e aumentem os recursos de *hardware* apenas quando há um aumento de suas necessidades;
3. A capacidade de pagar somente pelo que foi usado dos recursos computacionais, por exemplo, os processadores por hora e armazenamento por dia, e liberar recursos contratados facilmente quando não são mais necessários;

A computação em nuvem é composta por *Software* e *Hardware*, do ponto de vista de implementação, a computação em nuvem se destaca em 4 (quatro) vertentes [Mell and Grance 2009]:

- **Nuvem Privada:** oferece serviços para a própria organização, sendo operada e utilizada apenas pela mesma.
- **Nuvem Comunitária:** baseia-se em um ambiente de computação em nuvem compartilhado entre organizações com interesses em comum.
- **Nuvem Pública:** é um modelo que disponibiliza ambientes para o público em geral e são normalmente comercializadas por corporações com grande poder de armazenamento e processamento.
- **Nuvem Híbrida:** baseia-se na composição entre dois ou mais ambientes de estruturas distintas, nuvem privada e nuvem pública, por exemplo, gerando uma única nuvem, porém a conexão entre essas é feita a partir de tecnologias proprietárias.

De acordo com Mell e Grance (2009) [Mell and Grance 2009], computação em nuvem possui algumas características que se destacam, são elas:

- **Serviço sob demanda:** o provimento automatizado de funcionalidades computacionais, não necessitando de intervenção humana com o provedor do serviço;
- **Amplo acesso a serviços:** permite a disponibilização de recursos através da rede, habilitando o acesso a clientes de diferentes e diversos dispositivos que podem ser computadores, *smartphones*, dentre outros;

³<http://www.windowsazure.com/> Acessado em: 02/05/2016

⁴<http://www.dropbox.com> Acessado em: 02/05/2016

- **Multitenância:** permite o provimento de serviços a múltiplos usuários, e tais serviços podem ser alocados dinamicamente de acordo com a demanda;
- **Elasticidade:** correspondente à escalabilidade. Oferece ao usuário a sensação de ter os recursos disponíveis de forma ilimitada e a qualquer instante;
- **Tarifação:** segue o conceito *pay-as-you-go*, ou seja, o usuário paga somente pelo que for usado.

Dentre os principais desafios da computação em nuvem destacam-se: os desafios de prover segurança e a instável conexão de *Internet* em lugares geograficamente remotos. A segurança em ambientes em nuvem é constantemente contestada e se torna tema frequente de discursos científicos, questionando sempre sobre a segurança e o sigilo de dados armazenados na nuvem, principalmente em nuvens públicas. Locais geograficamente remotos apresentam baixas taxas de transferências de dados pela *Internet*, impossibilitando o uso eficiente de serviços em nuvem, em contraste de grandes centros com altas taxas de conexão com a *Internet*, possibilitando o uso de serviços em nuvem em sua total capacidade.

A utilização da computação em nuvens permite que os usuários possam alocar recursos dinamicamente, de acordo com sua necessidade. Entre esses recursos está o de armazenamento, o qual provê recursos e serviços de armazenamento baseados em servidores remotos que utilizam os princípios da computação em nuvem [Zeng et al. 2009]. Armazenamento em nuvem tem duas características básicas: a primeira trata da infraestrutura da nuvem, a qual baseia-se em *clusters* de servidores baratos; a segunda tem o objetivo de, através dos *clusters* de servidores, armazenamento distribuído e redundância de dados, fazer múltiplas cópias dos dados armazenados para alcançar dois requisitos: alta escalabilidade e alta usabilidade. A alta escalabilidade significa que o armazenamento em nuvem pode ser dimensionado para um grande aglomerado com centenas de nós ou *peers* de processamento. Alta usabilidade significa que o armazenamento em nuvem pode tolerar falhas de nós e que estas falhas não afetam todo o sistema [Deng et al. 2010].

Com a popularização da computação em nuvem, cresceu o número de sistemas utilizados para armazenamento em nuvem. O serviço oferecido por estes sistemas baseiam-se no *software* que é executado em um *cluster* de servidores, os quais armazenam em seus discos rígidos os arquivos dos clientes. Normalmente, cada cliente possui um processo que controla a transferência de arquivos entre a máquina do usuário e os servidores na nuvem. Este processo também tem a função de se certificar que os dados enviados sejam espalhados por outros servidores no *cluster* e manter a sincronia dos dados na máquina do cliente e os dados armazenados, ou seja, novos dados gerados na máquina do cliente deverão ser salvos na nuvem e, caso os dados locais sejam perdidos, recuperá-los [Machado 2013].

Com a expansão do mercado e das tecnologias em nuvem surgiram diversos sistemas de armazenamento de dados em nuvem, por exemplo, *DropBox*⁵, *Box*⁶, *JustCloud*⁷ e *Ustore*⁸, estes são alguns dentre os diversos sistemas de armazenamento em nuvem existentes [Rodrigues et al. 2015].

⁵<https://www.dropbox.com/> Acessado em: 02/05/2016

⁶<https://www.box.com/> Acessado em: 02/05/2016

⁷<http://www.justcloud.com/> Acessado em: 02/05/2016

⁸<http://usto.re/> Acessado em: 02/05/2016

Os sistemas de armazenamento em nuvem se mostram bastantes atrativos aos usuários, por fornecerem as opções de acesso, recuperação e armazenamento de arquivos de qualquer lugar e a qualquer hora. Desta forma, os usuários desses serviços podem armazenar arquivos em diversos formatos e tamanhos (.pdf, .doc, HTML, XML, RTF, arquivos compactados, imagens, arquivos de áudio e diversos outros formatos de arquivos), alguns desses serviços limitam o tamanho máximo de um único arquivo que pode ser armazenado, normalmente esta limitação acontece em sistemas públicos. Os sistemas de armazenamento em nuvem, em sua maioria baseiam-se nos princípios de *pay-as-you-go*, ou seja, o usuário paga pelo espaço que utilizar no sistema [Rodrigues et al. 2015].

As vantagens e atrativos apresentados por este tipo de sistemas atraem cada vez mais usuários, o que acaba contribuindo para o crescimento da massa de dados na nuvem. O que torna a atividade de filtragem de conteúdo considerado relevante, complexa e trabalhosa, fazendo com que usuários demandem mais tempo na busca por conteúdo de seu interesse. Observando este cenário é possível afirmar que a utilização de mecanismos de recomendação em sistemas de armazenamento em nuvem torna-se imprescindível, devido ao crescimento constante no volume de dados nesses sistemas [Rodrigues 2015].

3.2.1. Computação em Nuvem Como Negócio

A computação em nuvem traz para o mercado uma nova forma de ver e vender tecnologia. Essa mudança é semelhante ao que já é vivenciado pela população, como oferta de energia, água e serviços de telefonia. Essa característica da computação em nuvem está promovendo grandes mudanças nos modelos de negócios, desde pequenas empresas até às gigantes da computação. A cada dia, observamos o surgimento de empresas especializadas em fornecer serviços em nuvem. Esse cenário confirma de vez a computação em nuvem como tendência de mercado para o presente e para os próximos anos.

O setor público também caminha para essa tendência, o Serpro⁹, a Dataprev¹⁰ e outras instituições de serviços públicos estão se articulando para o provimento de serviços públicos de governo ligados às TICs de maneira compartilhada em nuvem.

A computação em nuvem oferece diversos nichos de negócios que podem gerar novos modelos de negócios, como segurança da informação, redução de custo com *hardware*, migração e compartilhamento de *softwares* legados e etc.

3.3 Modelos de Negócios

Cada empresa ou empreendimento possui um plano de negócio que é diferente de modelo de negócio. O modelo de negócio especifica o que é o negócio, em forma de um documento que descreva todas as características e atores envolvidos no negócio [Meira 2013]. A Tabela 2, resume a metodologia de criação de um modelo de negócio.

De acordo com Meira (2013) [Meira 2013], um bom modelo de negócio deve responder cinco perguntas principais sobre um negócio, são elas:

Quem paga? um modelo de negócio deve ter descrito quem serão os possíveis pagadores pelo produto/serviço que será desenvolvido.

⁹<https://www.serpro.gov.br/> Acessado em: 02/05/2016

¹⁰<http://portal.dataprev.gov.br/> Acessado em: 02/05/2016

TIPOLOGIA	ATRIBUTOS	PONTOS FORTES	PONTOS FRACOS
Definições Gerais	- Componentes que constituem o negócio. - Atributos gerais da indústria. - Metamodelo ou ontologia para modelo de negócios	As vantagens da agregação, ou seja, o entendimento das bases para a criação de valor da empresa.	A imagem transmitida se torna geral demais para passar qualquer coisa relevante sobre o negócio específico
Definições Amplas	- O método de fazer negócio. - Foco em todo o sistema do empreendimento. - A arquitetura para gerar valor. - Descrição de funções e relacionamentos.	- A criação de valor deve ser entendida através de toda a cadeia de valores da qual a empresa faz parte.	- Falta de suficiente foco nos principais processos de criação de valor. - Inclusão de fatores que não são completamente controlados pela empresa.
Definições Específicas	- Descrição do caráter único de aspectos internos. - Infraestrutura para geração de valor. - Considerações detalhadas de ligações, processos e rede de causa e efeito.	- O nível de detalhe respeita o funcionamento específico da firma. - Descrições precisas e relevantes.	- As considerações podem ficar muito específicas para fazer sentido. - Perda da Compreensão geral.

Figura 2. Modelo de Negócio [Meira 2013]

O que? outro ponto crucial de um modelo de negócio é descrever o que realmente é o produto/serviço que será entregue pelo empreendimento.

Para quem? deve ser explanado quem é o público alvo do negócio, por exemplo, donos de locadoras de automóveis.

Por quê? um bom modelo de negócio deve fundamentar e apresentar o porque será desenvolvido o produto esperado, a partir deste ponto, será possível delinear todos os outros pontos do negócio.

Como? neste ponto, deve ser apresentado como será desenvolvido o produto esperado como resultado.

3.3.1 Plano de Negócios

O plano de negócio é o detalhamento prescritivo de como o negócio deverá de ser feito, passo a passo, para se atingir o objetivo final do negócio. Um plano de negócio vai além do modelo de negócio, abrangendo também, modelos conceituais, planos de execução, plano financeiro, plano de pessoal e valores do negócio [Meira 2013]. A seguir na Figura 3 é apresentado um exemplo de plano de negócio.

PLANO DE NEGÓCIOS				
Análise do Ambiente				
Mercados Financeiros	Tamanho e crescimento do mercado	Necessidades do mercado		Competidores
Modelos Conceituais				
Modelo de Negócio			Proposta de Valor	
Planos de Execução				
Plano de Capital	Plano de tecnologia	Plano de Marketing		
		Plano de canais	Estratégia de precificação	Plano de vendas
Plano Financeiro				
Plano de Pessoal				
Valores Maximizados				
Valor para acionista			Valor para o cliente	

Figura 3. Plano de Negócio [Meira 2013]

O modelo de negócio é central e essencial para se escrever um plano de negócios. O Plano de negócio, detalhará toda a estrutura do negócio, apresentando passo a passo

cada etapa de desenvolvimento da startup. Em apresentações de negócios, um “Plano de Negócio” bem definido pode ser o diferencial de um negócio que receberá ou não um investimento.

Os modelos de negócio são formas de como conduzir uma organização para a qual ele foi desenvolvido. Especifica padrões, técnicas e ferramentas necessárias para atingir o objetivo do negócio. O modelo de negócio é a forma como uma organização cria, entrega e captura valor. Baseado nessa definição, surgiram modelos adaptados ao contexto das *startups*, denominados de modelos de gestão de negócios inovadores, esses têm substituído os tradicionais modelos e se adaptado às mudanças no cenário que têm como fatores críticos e fundamentais a agilidade e o dinamismo [Torres and de Souza 2015].

3.3.2 Modelo de Negócio Canvas

Quando novas empresas/negócios são criados (idealizados), têm-se a opção de escrever um modelo de negócios para explicar o novo empreendimento. Em geral, um modelo de negócio é um documento que descreve a visão para uma empresa e suas projeções financeiras. Normalmente conhecido por ser um documento longo e trabalhoso de ser desenvolvido, mas de suma importância para um novo negócio. Nesta seção, será apresentado o modelo de negócio *Canvas*, também conhecido como um forma resumida de modelo de negócio [Pesce 2012].

O *Business Model Canvas* (Modelo de Negócio Canvas), se popularizou com o grande crescimento dos projetos de startups, por ser um forma rápida, prática e simples de escrever um modelo de negócio. O *Canvas*, é sempre utilizado em eventos onde empreendedores possuem pouco tempo para idealizarem e apresentarem uma nova proposta de negócio inovador [Rodrigues 2015]. Na Figura 4, é apresentado o *Canvas*.

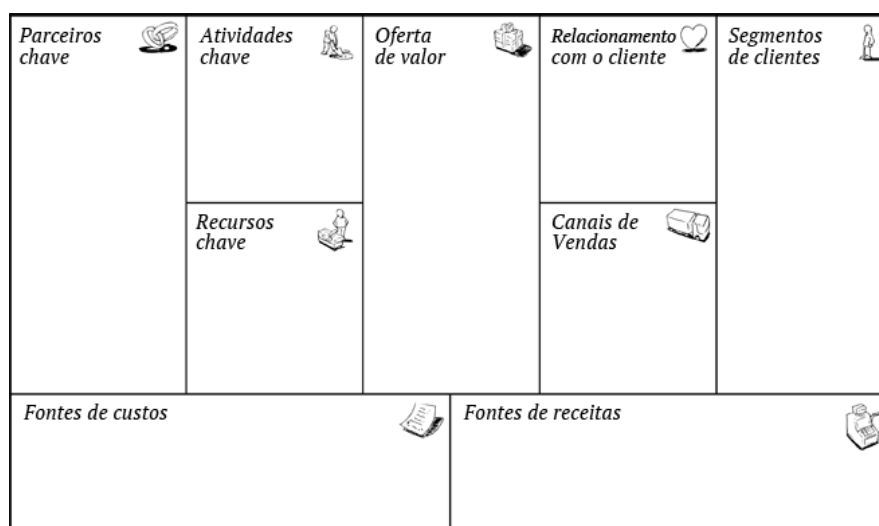


Figura 4. Modelo de Negócio Canvas

O *Canvas* é uma ferramenta de gerenciamento estratégico, um mapa dividido em blocos resumindo os principais itens de um modelo de negócios. O modelo foi proposto por Alexander Osterwalder¹¹, com o objetivo desimplificar o árduo trabalho de criar um

¹¹<http://www.businessmodelgeneration.com/canvas/bmc>

modelo de negócio, que na maioria das vezes se torna uma tarefa cansativa e complexa. O modelo *Canvas* é uma ferramenta ideal para novos empreendedores que estão iniciando seus *startups*. Cada bloco deve ser preenchido de forma objetiva e clara [Rodrigues 2013], a seguir são detalhados cada item do modelo *Canvas*.

Parceiros chave: neste item devem ser descritos as instituições ou pessoas que são parceiros da *startup*, contribuindo de forma direta ou indiretamente com o modelo de negócio da *startup*, por exemplo, a empresa que irá fornecer os servidores para hospedagem do site da *startup*.

Atividades chave: são as principais atividades que devem ser desenvolvidas pelo time da *startup*, para que seja produzido o produto objetivado no projeto, por exemplo, desenvolver o site da *startup*.

Recursos chave: são todos os recursos que o time da *startup* já possui para o desenvolvimento do projeto, por exemplo, mão de obra qualificada.

Oferta de valor: são os valores que serão agregados com o projeto, o diferencial do produto com os demais similares a ele, por exemplo, uma nova tecnologia como todos os serviços em nuvem.

Relacionamento com o cliente: são os meios em que a *startup* se relacionara com os clientes, por exemplo, atendimento via telefone.

Canais de venda: são os meios por onde o produto chegará ao cliente alvo, por exemplo, redes sociais.

Segmentos de clientes: neste item devem ser apresentados quais são os clientes alvo do produto, por exemplo, crianças e idosos ou pequenos empresários.

Fontes de custos: neste item são descritas todas as despesas do projeto, por exemplo, contratação de pessoal.

Fontes de receitas: neste item são descritas as fontes de receitas do projeto, de onde e como serão extraídas as receitas para a manutenção e o lucro da *startup*, por exemplo, venda de assinaturas.

Na Figura 5, é apresentado um exemplo do modelo de negócio *canvas* desenvolvido pela *usHouse*¹². O exemplo a seguir, é baseado na rede social “Facebook”.

A seguir são detalhados cada item do modelo *Canvas* montado para a rede social Facebook¹³:

Parceiros chave: neste item foi lista o “Bing” como parceiro chave da rede social, este é um sistema de busca *online*.

Atividades chave: foi determinado que o desenvolvimento e manutenção da plataforma é uma atividade chave para o funcionamento da rede social.

Recursos chave: como recursos chave foram citados quatro itens: equipe, marca, base de usuários e infraestrutura de tecnologia. Estes são recursos já disponíveis para o desenvolvimento e início da operação da rede social.

Oferta de valor: os principais valores apresentados pela rede social Facebook foram, a interação social *online* entre seus usuários e a grande audiência que proporcionada pelo Facebook.

Relacionamento com o cliente: um dos principais relacionamentos com o cliente foi definido como a própria interação na rede social.

¹²<http://ubhouse.com.br/montando-o-modelo-de-negocio-com-o-canvas-2/>

¹³<https://www.facebook.com/>

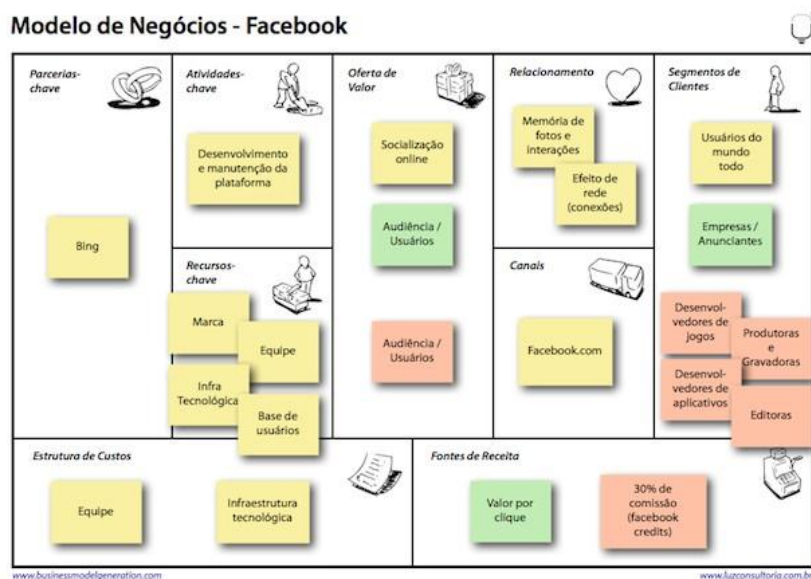


Figura 5. Modelo de Negócio Canvas da Rede Social Facebook

Canais de venda: o canal de entrega do produto ao cliente é a página da rede social.

Segmentos de clientes: o seguimento dos clientes é bem genérico e aberto, visto que qualquer pessoa que tiver conexão com a *Internet* pode se cadastrar e usar a redes social

Fontes de custos: as principais fontes de custos são a contratação de equipe capacitada e infraestrutura tecnológica.

Fontes de receitas: as fontes de receitas são providas de anúncios que pagam por cada clique que um determinado anúncio receber na rede social.

3.3.3 Modelos de Negócios Inovadores

Novos negócios inovadores não precisam ser novos no sentido de criar uma nova empresa ou um novo produto/serviço inédito no mercado. Um novo negócio poder algo desenvolvido em uma empresa que já se consolidou no mercado, o novo poder ser a reformulação de um conceito, ou até produtos que realmente saíram do zero. A noção de negócios inovadores descreve produtos/serviços que mudam comportamentos no mercado e no cotidiano das pessoas [Meira 2013].

A inovação resulta da busca por soluções diferenciadas e elegantes que visem resolver um problema real ou atender uma demanda latente, que gerem valor para os clientes e/ou que alcem a organização a uma posição privilegiada no mercado, onde elegância é encontrar a solução certa para um problema com simplicidade, criatividade, inteligência, sutileza, economia e qualidade.

A sobrevivência de uma organização está relacionada a como ela reinventa o setor e se diferencia dos concorrentes. A diferenciação, ganha lugar como uma das estratégias mais poderosa no mundo dos negócios e principalmente da inovação [Meira 2013].

De acordo com Meira [Meira 2013], inovação é um meio para se atingir diferenciação, e que, é a única fonte de vantagens competitivas sustentáveis. Essa

diferenciação, somente pode ser alcançada em boa parte na execução do negócio na prática, no ciclo de vida do negócio utilizando técnicas, métodos e processos.

3.4 Técnicas e Ferramentas

Nesta seção, são apresentadas algumas das principais técnicas, processos, métodos e ferramentas utilizadas por empreendedores de diversos níveis de negócios no desenvolvimento de modelos de negócios.

3.4.1 Equipe

Um *startup*, é um processo contruído por um time, que estará em constante processo de transformação. No decorrer do processo é normal que membros sejam considerados inadequados para o andamento do projeto. Desta forma, a equipe estará em constante construção e amadurecimento durante todo o projeto. Meira [Meira 2013] afirma que para se empreender é preciso preparo, dedicação, foco e autocrítica.

Meira [Meira 2013] sugere que o principal ativo de um *startup* é o seu capital humano, responsável por gerir e criar o produto final a partir de um modelo de negócio. Alguns pesquisadores sugerem que uma equipe ideal para iniciar um projeto de “*Startups de Software*”, deveria no mínimo ser composta por quatro membros: desenvolvedor, empreendedor, vendedor e um *design*.

- **Empreendedor:** o empreendedor é o responsável pela visão do negócio e por criar às melhores perspectivas ao time. O empreendedor se possível deve ser alguém inspirador, que passe confiança aos membros da equipe e ao mesmo tempo mantenha os pés no chão nas projeções do negócio.
- **Desenvolvedor:** em *startups de software* é essencial que a equipe conheça a tecnologia que será utilizada no desenvolvimento do produto final. A terceirização da tarefa de implementação, raramente é possível, justamente pelo fato de que uma *startup* ainda busca um modelo de negócio sólido e recursos financeiros. Desta forma, uma saída é convidar desenvolvedores para fazerem parte da equipe.
- **Design:** todo bom produto de *software* deve ser apresentável e proporcionar ao usuário um ambiente agradável. Diante deste cenário, é essencial que um novo produto tenha como membro um *design* capaz de proporcionar aos usuários novas experiências.
- **Vendedor:** é o responsável por apresentar o produto aos possíveis clientes, este membro será a primeira impressão que um possível cliente terá da *startup*. O vendedor será o responsável direto pelos ganhos financeiros e por atrair novos clientes, investidores e parceiros para o novo negócio.

3.4.2 Lean Startup

De acordo com Eric Ries em seu livro *The Lean Startup* (A Startup Enxuta) [Ries 2012], desenvolver um startup é um exercício de desenvolver uma instituição, portanto, envolve necessariamente administração. Muitas vezes, isso surge como uma grande surpresa para os aspirantes a empreendedores, pois suas associações com essas duas palavras são diametralmente opostas.

A metodologia “*Startup Enxuta*” baseia-se na experimentação e *feedback* dos clientes, com o objetivo de ajudar as Startups no desenvolvimento de produtos e serviços inovadores em uma relação estreita com os possíveis clientes (público alvo) [Torres and de Souza 2015]. O elemento central dessa metodologia é um processo espiral

composto das etapas: construir, medir e aprender.

No processo “*Startup Enxuta*”, conforme a Figura 6, a partir das idéias você constrói um produto mínimo viável (código), mede os resultados, coleta dados e aprende alguma lição. E continua a executar este laço de aprendizagem, o mais rápido possível, fazendo ajustes até atingir o casamento do produto com o mercado ou mudar algum item do modelo de negócios fazendo o pivô e começando tudo de novo. O objetivo é conseguir um modelo de negócio de valor, ou seja, que deixe o cliente feliz e gere lucro.

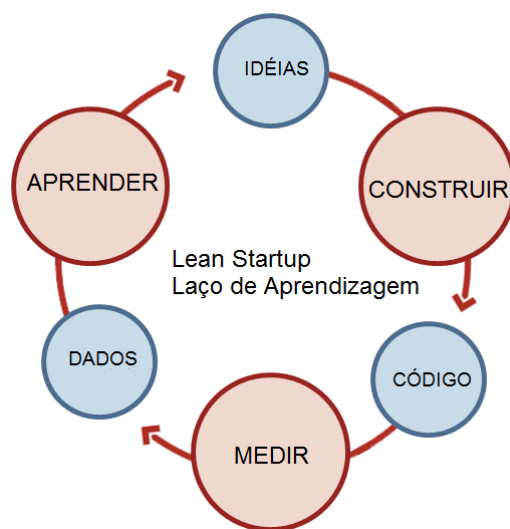


Figura 6. Processo “*The Lean Startup*” [Ries 2012]

- **Ideias:** segundo Eric Ries [Ries 2012] deve-se iniciar o processo através da definição da visão do negócio para identificar a estratégia, a visão é caracterizada como a descrição do negócio. A partir dessa definição, são montados os perfis dos possíveis clientes. Em seguida, a visão é decomposta nas hipóteses de valor e crescimento do negócio. As hipóteses de valor são formuladas para testar se o produto ou serviço de fato fornece valor aos usuários e, as hipóteses de crescimento têm a finalidade de verificar se o crescimento do produto está dentro das expectativas do negócio.
- **Construir e Código:** nessas fases é desenvolvido o MVP (Mínimo Produto Viável), que é um experimento que tem como objetivo validar as hipóteses. Este pode variar desde testes simples até protótipos.
- **Medir e Dados:** o produto das fases anteriores é testado na fase de medir por um período breve. O objetivo é garantir a efetividade dos testes realizados, nesta fase os dados são colhidos, analisados e avaliados para observar o progresso do negócio.
- **Aprender:** após a fase de medição é possível perceber que se deve perseverar a estratégia inicial ou realizar mudanças específicas na estratégia. Para ambas as hipóteses, a aprendizagem é que discerne para que hipótese deve-se seguir e, em seguida, o ciclo deve ser novamente executado, de forma mais ágil.

O proposta da *Lean Startup*, é fazer tudo da forma mais simples possível, usando o mínimo de velocidade para economizar dinheiro e diminuir riscos. Uma empresa enxuta, começa com um produto mínimo viável, através de um processo iterativo de aprendizagem e validação qualitativa, busca o ajuste do produto ao mercado para só então crescer em escala e estrutura [Rodrigues 2015].

3.4.3 MVP

Muitas *startups* tem ideias espetaculares para novos negócios e modelo de negócio bem definido, mas não conseguem entregar o produto idealizado de forma ágil ao mercado. Este cenário, acarreta em muitas vezes a perda de mercado, para grandes empresas que produzem o mesmo produto/serviço só que de forma ágil. *Startups* quase sempre não contam com grandes equipes e muito menos com recursos financeiros. Desta forma, uma boa alternativa é a produção de um produto mínimo viável [Rodrigues 2015].

Em startups, de acordo com Meira [Meira 2013], a primeira versão do produto deve ser um produto mínimo viável (MVP), ou seja, o produto mínimo concebível que pode encontrar um conjunto de clientes dispostos a pagar e utilizá-lo. Segundo Ries [Ries 2012], o MVP é uma das melhores técnicas para avaliar novos produtos/serviços no mercado.

O primeiro passo em criar um produto mínimo viável de sucesso, deve ser encontrar os usuários visionários ou evangelistas que querem e precisam do produto. Esses usuários serão capazes de ter a visão final do produto, por isso vão ignorar as falhas temporárias e acabarão ajudando a aprimorar o produto da Startup [Rodrigues 2013] [Ries 2012].

A ideia principal por trás do produto mínimo viável é, apresentar de forma ágil o produto idealizado e concebido como primeira versão, para que seja avaliado por usuários alvo do modelo de negócio *startup*. Desta forma, se o produto não atender as expectativas preliminares dos possíveis usuários, até mesmo o próprio modelo de negócio pode mudar de rumo e adaptar-se ao *feedback* dos usuários, economizando tempo e minimizando os riscos de ao final do projeto o produto resultante não ser aceito pelo mercado. Na Figura 7, é apresentado de forma simples o processo MVP.

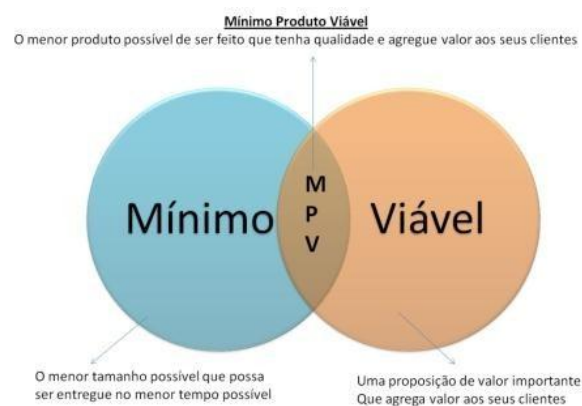


Figura 7. Processo MVP.

- **Mínimo:** é o menor tamanho possível, que possa ser entregue no menor tempo possível.
- **Viável:** uma proposta de valor que seja importante o suficiente para que seu principal cliente alvo utilize o seu produto.
- **Produto:** funcionalidades que apresentem o valor da proposta final, afim de que os usuários criem expectativas quanto à versão final do produto.

3.4.2 Oceano Azul

A estratégia do oceano azul visa prover uma nova forma para avaliar mercados, levantar as características que regem a concorrência, definir critérios de consumo, satisfação e

atendimento aos clientes, tendo como objetivo propor uma solução inovadora que se diferencie das demais [Rodrigues 2013].

A estratégia do oceano azul trata de dois cenários distintos. O primeiro é conhecido como “Oceano vermelho”, este possui vários consumidores sendo levados a consumir produtos e serviços de vários fornecedores que oferecem pouco ou nenhum diferencial entre seus produtos e serviços. Nesse cenário, podemos imaginar uma “sangrenta” competição por pequenos espaços de mercado entre os concorrentes em busca da sobrevivência.

O segundo cenário é o “oceano azul” cuja disputa por espaço não mais existe, pois neste caso os fornecedores oferecem algo diferenciado o bastante para tornar a concorrência irrelevante. O oceano azul é um novo espaço de mercado que ainda não havia sido descoberto, explorado pela concorrência.

O objetivo principal da estratégia “Oceano azul” é evitar as batalhas entre concorrentes e oferecer aos consumidores algo realmente novo, exclusivo e inovador, produzindo assim, a chamada inovação de valor, que alinha inovação com utilidade imediata e preço competitivo [Rodrigues 2013].

3.4.3 Modelo da Cauda Longa

O conceito de Cauda Longa está diretamente relacionado à economia da abundância, na qual o mercado apresenta diversos concorrentes que buscam oferecer produtos de massa aos consumidores. Já o mercado de oportunidades pouco exploradas ou inexistentes, este é mais conhecido na literatura como nicho. A cauda longa do mercado apresenta inúmeras variedades de nichos que muitas vezes ainda não foram exploradas pelo mercado atual. Além disso, existe abundância de consumidores na cauda longa, estes, por sua vez, nem mesmo sabem o que realmente precisam até que o produto ou serviço esteja disponível no mercado, exemplo disso, é o serviço do Netflix [Rodrigues 2013].

3.4.6 O que é um *pitch*?

Pitch é uma apresentação dos principais pontos que dura de 3 a 5 minutos com o objetivo de despertar o interesse dos ouvintes (clientes ou investidores) pelo seu negócio. O *pitch* deve conter apenas informações essenciais e que apresentem o diferencial do negócio, em regra é apresentado verbalmente com poucos *slides*. Em regra um *pitch* deve conter cinco tópicos genéricos, esses podem mudar de acordo com a peculiaridade de cada negócio:

1. Qual é a oportunidade.
2. O Mercado que irá atuar.
3. Qual é a sua solução.
4. Seus diferenciais.
5. O que está buscando.

O fundado dos Anos do Brasil¹⁶ apresenta uma maneira de como elaborar uma apresentação de um negócio baseada em cinco *slides*:

- **Slide 1 Identificando a Oportunidade:** o primeiro *slide* deve indicar qual a oportunidade que o negócio irá atender, qual o mercado e as necessidades que o mesmo tem e não são atendidas pelos negócios majoritários. Exemplo: Nós iremos resolver o problema das perdas na distribuição de água;

- **Slides 2 e 3 Apresentando a sua Solução:** esses slides devem apresentar a solução e a justificativa de porque ela irá resolver os anseios do mercado alvo;
- **Slide 4 Destacando seus Diferenciais:** nessa etapa devem ser apresentadas todos os diferenciais do negócios, do conjunto do empreendimento. Vale destacar também o diferencial da equipe e a evolução no desenvolvimento do produto. Exemplo: Nossa tecnologia, diferentemente do maior empreendedor deste mercado, não precisa que se instalem medidores específicos, pois monitoramos o fluxo de água por nosso equipamento de detecção;
- **Slide 5 Explanando sua Proposta:** aqui deve ser apresentado o estado atual do negócios, do desenvolvimento do produto e o cronograma de trabalho até a data em que o produto possa ser disponibilizado para os clientes. Exemplo: Já temos um protótipo funcional testado e avaliado pela companhia X e estamos buscando um investimento para completar o desenvolvimento, fabricar as unidades piloto e fechar os primeiros contratos.

3.5 Transformando Ideias em Negócios

A cada dia cresce o interesse pela combinação de empreendedorismo com inovação e tecnologia. Os estímulos para essa combinação geralmente envolvem o estabelecimento de ambientes especializados, como parques tecnológicos, incubadoras de empresas e aceleradoras de negócios. No Brasil, está acontecendo um crescimento expressivo no número de startups e aceleradoras de empresas [Ribeiro 2015].

De acordo com Ribeiro [Ribeiro 2015], o uso comum dos termos incubadoras e aceleradoras, muitas vezes são mal empregados. Incubadoras tem o objetivo de nutrir negócios nascentes amortecendo-os de seu ambiente, e proporcionando aos novos negócios um ambiente propício para o seu crescimento em um espaço protegido das forças de mercado. Aceleradoras tem o objetivo de aumentar a velocidade das interações com o mercado visando ajudar novos negócios a se adaptarem de forma rápida as suas exigências [Cohen 2014].

3.5.1 Incubação

O investimento público e privado em Incubadoras de Negócios nasce da demanda de mercado por competitividade e pela necessidade de inovação. O aumento de pequenas empresas e de novos empreendedores, é visto como fonte fundamental para a geração de empregos. Esse cenário se dá pelo fato de as incubadoras serem definidas como capazes de aumentar as chances de sobrevivência de empresas incubadas.

¹⁶<http://www.anjosdobrasil.net/>

O processo de incubação tem por objetivo alavancar talentos empreendedores e aumentar a velocidade do desenvolvimento de produtos inovadores, acelerando assim o desenvolvimento do negócio. Para atingir esse objetivo, existe uma provisão para empresas emergentes de espaço flexível para as empresas, equipamentos compartilhados e serviços administrativos, suporte em negócios, *marketing*, construção de equipes, obtenção de capital e consultoria de empreendedores experientes.

Chandra [Chandra 2009] define às incubadoras como uma instituição capaz de promover um céu de brigadeiro “*safe heaven*” ao agregar valor a negócios nascentes via serviços tangíveis e intangíveis. Os tangíveis podem ser escritórios com custo subsidiado, ambientes compartilhados e estrutura de equipamentos gerais compartilhada. Os intangíveis podem ser a presença de consultores com experiência em negócios ou no ramo específico do negócio. Dessa forma, agrega-se conhecimento e experiência aos novos empreendedores.

3.5.2 Aceleração

Aceleradoras de negócios são organizações formadas por empreendedores experientes que provêm serviços, espaços, mentoria, rede de contatos e conhecimentos em diversas áreas na criação de novos negócios. A assistência, geralmente é oferecida na construção do grupo de empreendedores, o ajuste da ideia e a mentoria sobre o modelo de negócio e como será o seu lançamento no mercado. Após essas contribuições, as melhores ideias são selecionadas para uma apresentação para investidores, potenciais parceiros e clientes. O objetivo das aceleradoras é de oferecer o suporte necessário para as empresas desenvolverem seu modelo de negócios e alinharem-se com melhor eficiência em seus devidos mercados [Lynn 2012].

A seguir apresentamos uma estrutura básica para aceleradoras composta por cinco características principais:

1. **Um processo de aplicação altamente competitivo:** os programas, geralmente, têm inscrição *online*, o que permite que a aplicação seja em nível internacional e a seleção é feita por especialista com experiência no mercado.
2. **Provisionamento financeiro para sustentar a operação:** o investimento oferecido está geralmente relacionado ao custo para que cada empreendedor consiga se sustentar durante o período de aceleração em dedicação exclusiva ao projeto.
3. **Foco em pequenos times, e não em indivíduos:** geralmente, programas de aceleração não selecionam empreendedores sozinhos, pois considera-se que o trabalho necessário para levantar um startup exija mais de uma pessoa. Nem sempre, o time é selecionado pela ideia apresentada, mas sim, pelas características dos membros da equipe.
4. **Tempo limitado de suporte oferecido e mentoria intensiva:** o tempo limitado é relativo ao lançamento de um produto inicial (MVP) no mercado, mas também está ligado ao objetivo de criar um ambiente de alta pressão voltado para rápido progresso e aprendizado.
5. **Projetos acelerados em lotes:** a cada período de aceleração, entra um novo grupo de times empreendedores. A vantagem principal é o suporte gerado de maneira colaborativa entre as empresas, onde empreendedores de distintas áreas ajudam seus pares a resolverem problemas com competências peculiares em forma de reciprocidade indireta.

De acordo com Ribeiro [Ribeiro 2015], existem cinco perfis de aceleradoras no Brasil:

1. **Aceleradoras Abertas:** oferecem investimentos financeiros com contrapartida ligada à propriedade na empresa. Voltadas para empresas de alto crescimento e geralmente centradas em tecnologias de comunicações e informação. Exemplo: *Acceleratech*.
2. **Aceleradoras Corporativas:** oferecem investimento financeiro com contrapartida ligada à propriedade da empresa. Concentram-se nas áreas de atuação das empresas mantenedoras dos programas. Exemplo: Wayra (Telefônica).
3. **Pré-aceleradora:** não possui investimento financeiro, a duração é reduzida (em média 5 semanas), operação financeira do modelo dependente de patrocínios, eventos e parcerias diversas. Exemplo: *Startup Farm*.
4. **Aceleradoras Sociais:** programas centrados no desenvolvimento de negócios sociais, comumente sem relação financeira estabelecida. Operação geralmente vinculada a fundos de investimentos direcionados a esse perfil. Exemplo: Instituto Quintessa.
5. **Aceleradoras Públicas:** ofertam investimentos financeiros sem contrapartida. Possuem estímulo de órgãos públicos com interesse em desenvolvimento econômico regional. Exemplo: SEED do Governo de Minas Gerais.

3.5.3 Investidores Anjo

O investidor anjo geralmente é uma pessoa física com experiência (ex-empresário ou empreendedor) e possui recursos suficientes para alocar uma parte em novos negócios. Ser investidor anjo não é uma atividade filantrópica com fins sociais. Os investidores anjos aplicam seus recursos em negócios promissores e com grande potencial de retorno, consequentemente esses negócios terão impacto na sociedade. O termo anjo é utilizado pelo fato do investidor não ficar limitado apenas a parte financeira, eles podem contribuir com experiência, conhecimentos e sua rede de contatos para orientar e buscar novos direcionamentos para o negócio. Principais características de investimento-anjo:

1. É efetuado por profissionais experientes, que agregam valor para o negócio com seus conhecimentos, experiência, redes de contatos e recursos financeiros;
2. Adquiri normalmente uma participação no negócio e permanecendo como sócio minoritário;
3. Não possui posição executiva na empresa, mas apoia os empreendedores na forma de conselheiro (mentor).

3.5.4 Marketing Digital

Tirar uma ideia do papel e transformá-la em uma *startup* promissora é uma das atitudes mais difíceis para os empreendedores criativos. Após as etapas de edificação do projeto e captação de investidores, chega a hora de definir quais estratégias de marketing devem ser seguidas, e as digitais são uma boa escolha.

Na etapa de *marketing* é importante buscar otimização e campanhas com *links* patrocinados. O objetivo é fazer com que a *startup* apareça nas primeiras posições do *Google* quando um usuário pesquisa por algum termo relacionado ao segmento em que ela

atua. Outro importante fator é alimentar os usuários com conteúdo de qualidade. Vivemos na era da informação, dessa forma, é crucial para um negócio gerar e disponibilizar conteúdo para seus usuário, isso pode ser alcançado com páginas da *startup* nas principais redes sociais utilizadas pelos possíveis usuários.

3.6 Tendências

Nesta seção são apresentadas 10 tendências por Silvio Meira ¹⁷. São apresentadas tendências que já eram evidentes em 2015 e que podem se consolidar no ano de 2016 ou em um futuro não muito distante. O impacto dessa consolidação pode ser relevante para pessoas, instituições, organizações, economia, sociedade e para o mercado em rede, que já é uma realidade.

1. **A web se tornará cada vez mais fragmentada**, face a legislações nacionais que exigem serviços de exclusividade local. Vários fatores contribui para isso, como segurança de informações nacionais;
2. **Teremos mais conflitos tecno-legais**, um exemplo desses conflitos é o caso recente do bloqueio dos serviços de *Whatsapp* pela justiça brasileira. A tendência é que esses conflitos aconteçam com maior frequência, principalmente entre as gigantes da telecomunicação, usuários e governo;
3. **Startups terão mais dificuldade de entrar nos mercados globais**, como consequência das duas primeiras tendências citadas. Será mais comum que surja um mercado de oportunidades no Brasil, ao mesmo tempo em que outros mercados ficaram mais distantes, principalmente devido à variação cambial e as incertezas políticas vividas no Brasil;
4. **Uma rede de segunda classe**, É provável que surja no Brasil um nicho de mercado que não se encaixe em outras economias, isso pode acontecer devido as diferenças de investimentos tanto do setor privado quanto público;
5. **Internet das coisas**, a década atual é onde está sendo estabelecidas as fundações, os principais componentes e protocolos para que as coisas sejam informatizadas;
6. **Smart cities**, são um caso particular de ambientes inteligentes, e a cada dia as cidades estão cada vez mais inteligentes ou, pelo menos, espertas;
7. **Montanhas de dados** ou, simplesmente *Big Data*, com o crescimento do uso do meio digital montanhas de dados estão se formando em mais diferentes ambientes;
8. **Drones e robôs** chegaram para ficar. Os drones que estão por aí não passam de aeromodelos de última geração e controlados por um piloto, mas já estão chegando os drones que são autônomos, principalmente em ambientes militares e comerciais. Robôs já chegaram a mais tempo, mas nessa década emplacaram de vez. Em um futuro breve, será normal a presença de robôes e drones no cotidiano das pessoas;
9. **Inteligência Artificial**, parte do que robôs podem fazer depende de Inteligência Artificial (IA). IA é outra revolução digital que está praticamente pronta para ser usada, em larga escala. IA ganhou mais relevância com o crescimento de projetos com *big data* e a proliferação de recursos robóticos;
10. **Insegurança de informação será a norma e não exceção**, essa tendência está diretamente relacionada com todas as tendências apresentadas aqui, o principal fator

¹⁷<http://boletim.de/silvio/tendencias-para-2016-e-depois/>

que contribui para esse cenário é a dependência crescente de dados, informações e conectividades dos usuários;

3.7 Considerações Finais

Este artigo apresentou uma visão sobre computação em nuvem, modelos de negócios, *startups* e algumas das principais técnicas usadas por empreendedores, como Modelo de Negócio *Canvas*, *Startup Enxuta*, Estratégia do Oceano Azul, Mínimo Produto Viável (MVP) e o Modelo *Calda Longa*. Foram apresentados conceitos sobre inovação empreendedora e a evolução do empreendedorismo na era da computação em nuvem. No universo do empreendedorismo existem inúmeras técnicas de trabalho e otimização para o desenvolvimento de negócios, e neste trabalho apresentamos apenas algumas das mais difundidas no meio das *startups*.

Referências

- Chandra, A.; Fealey, T. (2009). Business incubation in the united states, china and brazil: A comparison of role of government, incubator funding and financial services. *International Journal of Entrepreneurship*, 13:67 – 86.
- Cohen, S. Hochberg, Y. V. (2014). *Accelerating Startups: The seed accelerator phenomenon*. Available at SSRN 2418000.
- Deng, J., Hu, J., Liu, A., and Wu, J. (2010). Research and application of cloud storage. In *Intelligent Systems and Applications (ISA), 2010 2nd International Workshop on*, pages 1–5.
- Lenk, A., Klems, M., Nimis, J., Tai, S., and Sandholm, T. (2009). What’s inside the Cloud? An architectural map of the Cloud landscape. In *Software Engineering Challenges of Cloud Computing, 2009. CLOUD.ICSE Workshop on*, volume 0 of *CLOUD ’09*, pages 23–31, Washington, DC, USA. IEEE.
- Lynn, H. D.; Radojevich-Kelley, N. (2012). Analysis of accelerator companies: An exploratory case study of their programs, processes, and early results. *Small Business Institute Journal*, 8:54 – 70.
- Machado, M. A. S. (2013). Uma abordagem para indexação e buscas full-text baseadas em conteúdo em sistemas de armazenamento em nuvem. Master’s thesis, Universidade Federal de Pernambuco (UFPE).
- Meira, S. L. (2013). *Novos Negocios Inovadores de Crescimento Empreendedor no Brasil*. ISBN 978-85-77344123.
- Mell, P. and Grance, T. (2009). The NIST Definition of Cloud Computing. Technical report, National Institute of Standards and Technology, Information Technology Laboratory.
- Pesce, B. (2012). *A menina do vale: como o empreendedorismo muda a sua vida*. ISBN 978-85-7734-280-8.
- Pressman, R. S. (2011). *Engenharia de Software: Uma Abordagem Profissional*. ISBN 978-85-63308-33-7.

- Ribeiro, Artur Tavares Vilas Boas, P. G. A. O. L. M. (2015). Um fim, dois meios: Aceleradoras e incubadoras no Brasil. *Anais do XVI Congresso Latino-Iberoamericano de Gestão da Tecnologia*.
- Ries, E. (2012). *A Startup Enxuta*. ISBN 978-85-8178-004-7.
- Rodrigues, R. B., da Silva, C. M. R., Durao, F. A., Assad, R. E., Garcia, V. C., and Meira, S. R. L. (2015). A file recommendation model for cloud storage systems. In *Proceedings of the Annual Conference on Brazilian Symposium on Information Systems: Information Systems: A Computer Socio-Technical Perspective - Volume 1*, SBSI 2015, pages 16:111–16:118, Porto Alegre, Brazil, Brazil. Brazilian Computer Society.
- Rodrigues, Ricardo Batista, R. T. A. d. O. e. R. R. d. S. (2013). Startups dirigidas a inovação de software: Da universidade ao mercado. *Anais da III Escola Regional de Informática de Pernambuco (ERIPÉ)*, 2:162 – 169.
- Rodrigues, Ricardo Batista, S. C. M. G. V. C. M. S. R. L. (2015). Criando startups: métodos, processos e técnicas e ferramentas. *Anais do XI Simpósio Brasileiro de Sistemas de Informação*, 4:13–17.
- Torres, N. N. d. J. and de Souza, C. R. B. (2015). Software startup ecosystems: Initial results in the state of para. In *Proceedings of the Annual Conference on Brazilian Symposium on Information Systems: Information Systems: A Computer Socio-Technical Perspective - Volume 1*, SBSI 2015, pages 12:83–12:86, Porto Alegre, Brazil, Brazil. Brazilian Computer Society.
- Vaquero, Luis M., R.-M. L., Caceres, J., and Lindner, M. (2008). A break in the clouds: towards a cloud definition. *SIGCOMM Comput. Commun. Rev.*, 39(1):50–55.
- Vogels, W. (2008). A head in the clouds - the power of infrastructure as a service. *First workshop on Cloud Computing and in Applications (CCA 2008)*.
- Zeng, W., Zhao, Y., Ou, K., and Song, W. (2009). Research on cloud storage architecture and key technologies. In *Proceedings of the 2nd International Conference on Interaction Sciences: Information Technology, Culture and Human*, ICIS '09, pages 1044–1048, New York, NY, USA. ACM.

Capítulo

4

Sistemas LBS, Internet das Coisas e Computação Vestível: Usando a Computação Sensível ao Contexto para Desenvolver as Aplicações do Séc. XXI

Marcio Pereira de Sá

Abstract

Ubiquitous Computing allows use of computational resources and access to information anywhere and anytime. It benefits greatly of Mobile Computing, through the use of mobile devices and applications and Context-Aware Computing which deals with the research and development of applications that understand e react to changes in the environment in which they execute. Thus, this chapter aims to make an introduction to Mobile and Ubiquitous Computing, analyzing its history and evolution, its main concepts and definitions, as well as new types of applications and computer systems can be created, such as Location-Based Systems, Wearable Computing and Internet of Things.

Resumo

A Computação Ubíqua permite o uso dos recursos computacionais e o acesso à informação a qualquer momento e em qualquer lugar. Ela se beneficia muito da Computação Móvel, através do uso dos dispositivos e aplicações móveis e da Computação Sensível ao Contexto que se ocupa da pesquisa e do desenvolvimento de aplicações que percebam e reajam às mudanças no ambiente em que estão sendo executadas. Assim, este capítulo tem por objetivo fazer uma introdução à Computação Móvel e Ubíqua, analisando sua história e evolução, seus principais conceitos e definições, bem como os novos tipos de aplicações e sistemas computacionais que podem ser criados, como os Sistemas Baseados em Localização, a Computação Vestível e a Internet das Coisas.

4.1. Introdução

A cada dia, nota-se que os dispositivos computacionais tem ganhado cada vez mais espaço na vida cotidiana de grande parte da população mundial. O uso desses dispositivos eletrônicos vai desde o entretenimento (jogos, áudio e vídeo), passando por assistentes importantes no trabalho e nos estudos, até exercerem o papel essencial em muitos setores de nossas sociedades contemporâneas, como no controle dos tráfegos aéreo e urbano e no gerenciamento de aparelhos em UTIs que ajudam a manter pacientes vivos. Deve-se ressaltar também a função importante que eles exercem no gerenciamento do sistema elétrico e até de usinas nucleares.

Como se vê, a humanidade é atualmente muito dependente desses sistemas computacionais e por isso, o estudo, a pesquisa e o desenvolvimento de sistemas controlados por computador têm sido tão importantes para governos e cidadãos em todo o planeta.

Desde a década de 1940, com a criação dos primeiros computadores eletrônicos, se tem visto uma grande evolução tanto no hardware quanto no software disponibilizados nestes dispositivos. Mas foi a partir da década de 1980 que uma importante área da Computação, conhecida como Computação Ubíqua começa a sua jornada [Weiser 1991]. Unindo-se à Computação Móvel, outra área também muito importante atualmente, foi possível chegar à criação de diversas tecnologias e subáreas de pesquisa e desenvolvimento, como a Computação Vestível (*Wearable Computing*), os Sistemas Baseados em Localização (LBS – *Location-based System*) e a Internet das Coisas (IoT – *Internet Of Things*).

Todas essas áreas se beneficiaram e ainda se beneficiam muito dos conceitos e aplicações de um outro campo de estudos, conhecido como Computação Sensível ao Contexto. Este campo se preocupa com o desenvolvimento de aplicações que percebam e reajam a mudanças do ambiente no qual estão inseridas como as mudanças na localização dos usuários, nos recursos computacionais dos dispositivos, na luminosidade e temperatura do local onde o usuário se encontra. Todas essas e muitas outras informações são denominadas *informações de contexto* e são providas, direta ou indiretamente, por sensores.

Este trabalho está organizado da seguinte forma: a seção 4.2 analisa a Computação Móvel, sua história e evolução, as diferentes gerações da telefonia móvel, bem como as principais categorias de dispositivos móveis. Um estudo sobre a Computação Ubíqua é realizado na seção 4.3, onde se discute a Computação Sensível ao Contexto, as definições de contexto e suas classificações principais, além de analisar algumas aplicações importantes da Computação Ubíqua, como os Sistemas Baseados em Localização, a Computação Vestível e a Internet das Coisas. A seção 4.4 apresenta as considerações finais sobre o assunto.

4.2 Computação Móvel

A Computação Móvel é um termo usado para indicar o campo de pesquisa e desenvolvimento de dispositivos e aplicações cuja mobilidade é a principal característica e diferença em relação à computação tradicional. É, nos dias de hoje, uma das áreas da Tecnologia da Informação e Comunicação (TIC) mais importantes, pois tem

experimentado, nas últimas três décadas, um crescimento espetacular a ponto de alguns a considerarem como a *Quarta Revolução da Computação* [Loureiro 1998], cujos antecessores foram os grandes centros de processamento de dados da década de 1960, o aparecimento dos terminais de acesso nos anos 1970 e a popularização dos computadores pessoais e das redes de computadores na década de 1980, conforme ilustrado na Figura 4.1.

Como a Computação Móvel se ocupa com as pesquisas e o desenvolvimento tanto de dispositivos móveis (*hardware*), quanto de aplicações e sistemas que são executados nesses dispositivos (*software*), é importante analisar aspectos tecnológicos, tanto de hardware quanto de software, quando se quer aprofundar no estudo desta área da Computação.

Assim, em relação ao hardware, algumas características são especialmente importantes para este estudo, como as várias gerações tecnológicas da telefonia móvel, que, de certo modo, ajudaram a popularizar os dispositivos móveis tão importantes no mundo atual. Quanto ao software, alguns aspectos importantes incluem os sistemas operacionais para dispositivos móveis, como o iOS (iPhone e iPad), o Android e Windows Phone/10, bem como as suas plataformas de desenvolvimento, *middleware* e *frameworks* de apoio.

Por serem muito importantes, esses e outros aspectos serão discutidos e analisados em maior profundidade nas subseções seguintes.

1ª Revolução da Computação	2ª Revolução da Computação	3ª Revolução da Computação	4ª Revolução da Computação
Grandes Centros de Processamento de Dados	Terminais de Acesso	Computadores Pessoais e Redes de Computadores	Computação Móvel
(década de 1960)	(década de 1970)	(década de 1980)	(décadas de 1990)

Figura 4.1 – As quatro grandes revoluções na Computação

4.2.1. Fatos Históricos Importantes da Computação Móvel

A Computação Móvel, segundo [Loureiro 1998], tem, como precursores, os sistemas analógicos de comunicação cuja origem está relacionada à descoberta de que a corrente elétrica produz um campo elétrico, atribuída a Hans Christian Oersted, em 1820.

Essa afirmação permitiu iniciar estudos e pesquisas que culminariam com o desenvolvimento de vários sistemas de comunicação, cujo primeiro foi o *telégrafo*. Em linhas gerais, este sistema, inicialmente baseado na comunicação com fio, ainda na metade do século XIX, permitia a transferência de palavras a longa distância através do código Morse. Mas, a partir do estudo da propagação de ondas eletromagnéticas, o que permitiu a descoberta das equações de Maxwell e os experimentos de Heinrich Hertz, foi possível a descoberta da radiotelegrafia por Guglielmo Marconi, no final do século XIX.

Como resultado dessas descobertas, em 1901, ondas eletromagnéticas de sinais de rádio atravessaram o Oceano Atlântico e, com isso, estavam lançados os fundamentos para a comunicação sem fio, tão importante e essencial no mundo atual.

Ainda de acordo com [Loureiro 1998], a segunda geração de sistemas de comunicação foi representada pelo *telefone*, inventado por Alexander Graham Bell, na década de 1870. Pode-se dizer que a terceira geração dos sistemas de comunicação se caracteriza pelo surgimento dos computadores eletrônicos, ainda na década de 1940. A grande contribuição dos sistemas computacionais neste campo foi a viabilização da comutação telefônica digital, permitindo o aumento exponencial de ligações telefônicas simultâneas, algo que seria inviável se feito apenas com operadores humanos (telefonistas).

É importante notar que, até este momento, a comunicação com fio, caracterizada pelo elevado custo de acesso remoto, ainda era predominante. Este cenário tornou os sistemas sem fio atraentes para as companhias telefônicas. Desse modo, várias pesquisas começaram, visando desenvolver novos sistemas de comunicação sem fio mais eficientes.

Historicamente, o primeiro sistema de comunicação móvel de uso prático foi um sistema de rádio utilizado, em 1928, pela polícia de Detroit, nos Estados Unidos. A partir da década de 1940, especificamente em 1947, a AT&T desenvolveu o IMTS (*Improved Mobile Telephone Service*), um sistema de transmissão, precursor dos sistemas celulares atuais, em que era necessário somente uma única torre de alta potência para atender à população de uma grande área ou cidade.

O primeiro sistema celular de uso amplo, entrou em uso nos Anos 1970, com o lançamento do AMPS (*Advanced Mobile Phone System*), pela AT&T. Seu uso inicial era em automóveis, atendendo a uma pequena quantidade de usuários, oferecendo uma capacidade de tráfego de informações ainda muito limitada. Em 1979, o Japão recebeu a primeira rede de telefonia celular no mundo. A segunda geração do sistema AMPS é lançada em 1983, permitindo o surgimento da primeira rede celular americana que inicialmente operava em Chicago e Baltimore [Loureiro 1998].

Muitas inovações no campo da telefonia celular aconteceram na década de 1990. Ainda no início da década, em 1991, ocorreu, nos Estados Unidos, a validação inicial dos padrões TDMA (*Time Division Multiple Access*) e CDMA (*Code Division Multiple Access*) e a introdução da tecnologia microcelular. Em 1992, surge o sistema Pan-Europeu GSM (*Groupe Spéciale Mobile – também conhecido como Global System for Mobile communications*). Em seguida, vemos a introdução do sistema CDPD (*Cellular Digital Packet Data*) e o início dos serviços PCS (*Personal Communication Services*), em 1994. Os projetos para a cobertura terrestre por satélites de baixa órbita, como o projeto Iridium, se iniciam em 1995.

A década de 2000 tem como fatos importantes a popularização do 3G e das redes Wi-Fi. Nesta década atual, notam-se vários fatos interessantes, como a disseminação das tecnologias 4G e o início dos estudos dos padrões 5G, além do aparecimento de novos

tipos de dispositivos, como os dispositivos vestíveis e o avanço nas pesquisas do que se está convencendo chamar de Internet das Coisas.

4.2.2. A Evolução da Telefonia Móvel e Suas Gerações

A Computação Móvel atual foi influenciada diretamente, durante as três últimas décadas, pela evolução da telefonia móvel. Por isso, é importante estudar a história da telefonia móvel para compreender melhor a evolução e o estado atual da própria Computação Móvel.

Neste sentido, para facilitar o estudo da história e a evolução deste área, é útil organizar seus fatos e acontecimentos em gerações. Cada geração se caracteriza especialmente pelas suas tecnologias de comunicação [Loureiro 1998]. As gerações e suas características principais são descritas e analisadas a seguir:

- *Primeira Geração - 1G*

A Primeira Geração da Telefonia Móvel foi caracterizada pelo uso de tecnologias analógicas, cujas consequências aparentes foram o emprego de aparelhos grandes, volumosos e pesados se comparados com os aparelhos das gerações subsequentes. Além disso, os principais padrões de comunicação dessa geração foram: o AMPS (*Advanced Mobile Phone System*), desenvolvido, ainda na década de 1970, nos Estados Unidos. Este é considerado o primeiro padrão de rede celular, de fato. O TAC (*Total Access Communication system*) foi uma contrapartida europeia ao AMPS. Mais tarde, foi desenvolvido uma evolução do TAC, denominada ETACS (*Extended Total Access Communication System*), porém todos esses sistemas de comunicação ainda eram analógicos.

- *Segunda Geração - 2G*

A passagem da tecnologia analógica para a digital deu origem à Segunda Geração da Telefonia Móvel. Esta mudança permitiu a diminuição tanto do peso quanto do tamanho dos aparelhos telefônicos móveis (celulares). Além disso, houve um aumento na eficiência do consumo de energia, permitindo aumentar a duração da carga das baterias destes aparelhos.

Os principais padrões de comunicação dessa geração são o CDMA (*Code Division Multiple Access*), o TDMA (*Time Division Multiple Access*) e o GSM (*Global System for Mobile communications*).

Além da voz, as tecnologias 2G incluíram a capacidade de transmitir pequenos volumes de dados digitais, como mensagens de texto curtas (SMS - *Short Message Service*) ou pequenas mensagens multimídia (MMS - *Multimedia Message Service*). Deve-se notar que o padrão GSM original permitia a transferência de dados a, no máximo, 9,6 Kbps.

- *Gerações Intermediárias: 2.5G e 2.75G*

Como em muitos países não houve uma migração imediata para a Terceira Geração (3G), foi necessário oferecer tecnologias intermediárias aos usuários do sistema de telefonia móvel desses países. Desse modo, foram disponibilizadas tecnologias denominadas Geração 2.5G (GPRS) e Geração 2.75G (EDGE).

O GPRS (*General Packet Radio System*) é, na verdade, uma extensão do padrão GSM, permitindo, teoricamente, velocidades de transferência de dados de até 114 Kbps. Por ser considerado uma evolução do padrão GSM de segunda geração, o GPRS foi chamado de Geração 2.5G.

Por outro lado, o padrão EDGE (*Enhanced Data Rates for GSM Evolution*) permitiu alcançar velocidades de transferência teóricas de aproximadamente 384 Kbps, viabilizando a utilização de aplicações multimídias mais simples. O padrão EDGE ficou conhecido como um padrão 2.75G.

- *Terceira Geração - 3G*

A Terceira Geração da Telefonia Móvel (3G) trouxe contribuições importantes tanto às operadoras de telefonia móvel quanto aos usuários desses serviços. No início da década de 2000, o IMT-2000 (*International Mobile Telecommunications for the year 2000*) da União Internacional das Comunicações (UIT) definiu e publicou o padrão 3G em substituição às tecnologias anteriores de Segunda Geração.

As principais características da Terceira Geração de telefonia móvel (3G) foram a possibilidade de transmissão de dados em altas taxas em comparação com a geração anterior, a compatibilidade mundial entre dispositivos e sistemas de telefonia 3G e ainda a compatibilidade desses serviços de terceira geração com os sistemas de segunda geração.

O UMTS (*Universal Mobile Telecommunications System*) foi considerado o principal padrão 3G. Este padrão utiliza uma banda de frequência de 5 MHz, o que permite alcançar velocidades de transferência de dados entre 384 Kbps e 2 Mbps. O padrão UMTS emprega a codificação W-CDMA (*Wideband Code Division Multiple Access*).

Em seguida, foi disponibilizada a tecnologia HSDPA (*High-Speed Downlink Packet Access*). Esta, por ser uma evolução do padrão 3G, é considerada um padrão superior, conhecido também por 3.5G. O HSDPA permite taxas de transferência entre 8 e 10 Mbps, o que possibilitou o desenvolvimento de novas aplicações móveis que fazem uso intensivo de transferência de dados.

- *Quarta Geração - 4G*

As redes de telefonia móvel de quarta geração (4G) estão baseadas em duas tecnologias/padrões principais: o LTE (*Long-Term Evolution*) e o Wi-Max (*Worldwide Interoperability for Microwave Access*, também conhecido como IEEE 802.16).

Com certeza, o aumento da velocidade de transmissão é a principal evolução dos padrões de quarta geração em relação aos padrões anteriores. Como exemplo, o padrão LTE, que tem se tornado o mais popular nos últimos anos, pode, em teoria, permitir que *downloads* sejam realizados a até 100 Mbps enquanto os *uploads* podem alcançar a taxa de 50 Mbps. Isto representa um avanço considerável na velocidade de transferência de dados se comparado ao padrão 3G original.

- *Quinta Geração - 5G*

Com o 4G se popularizando, as pesquisas para a próxima geração das redes móveis de comunicação já foram iniciadas e o início das operações dessas redes está previsto para 2020.

Em pesquisas realizadas com as tecnologias de quinta geração (5G), já foram atingidas velocidades de conexão de até 1 Terabyte por segundo. Esse aumento nas velocidades de transferência de dados viabiliza a disseminação e a popularização do uso de tecnologias como os dispositivos vestíveis e a Internet das Coisas (IoT), permitindo a criação do ambiente necessário para a consolidação da Computação Ubíqua, idealizada por Mark Weiser ainda no início da década de 1990 [Weiser 1991].

A Figura 4.2 ilustra alguns dos principais fatos sobre a evolução da Comunicação Móvel.

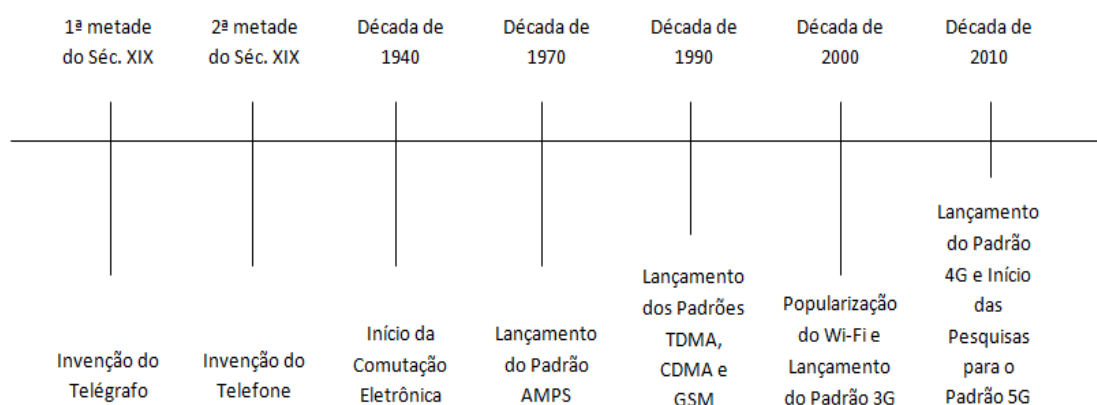


Figura 4.2 – Fatos importantes sobre a evolução da Comunicação Móvel

4.2.3. Classificação dos Dispositivos Móveis

É importante observar a grande quantidade de dispositivos móveis disponíveis no mercado atualmente ou que eram muito comuns até pouco tempo atrás. Esses aparelhos se subdividem em várias categorias ou tipos, dentre elas, destacam-se *telefones celulares convencionais*, *feature phones*, *smartphones*, *tablets*, *phablets* e *dispositivos vestíveis*.

Os *telefones celulares convencionais* foram os dispositivos móveis mais comuns nas duas últimas décadas e suas características principais são o baixo preço e a ausência de sistemas operacionais. São aparelhos muito simples com poucas funcionalidades além da possibilidade de realizar e receber chamadas telefônicas ou enviar e receber mensagens de texto (SMS). Possuem normalmente telas pequenas e são muito limitados em relação à quantidade de memória e armazenamento interno, ao poder de processamento e à forma de conectividade a redes sem fio e ao acesso à internet.

Os *feature phones* são dispositivos que oferecem, a exemplo dos telefones celulares convencionais, poucas funcionalidades, porém estes aparelhos já embutem melhorias nos recursos de conexão à internet e às redes sem fio que facilitam o uso destes para acesso a redes sociais, contas de email, etc. Tanto os *feature phones* como os

telefones celulares convencionais estão, cada vez mais, diminuindo sua participação no mercado global de dispositivos móveis, perdendo suas fatias de mercado para os *smartphones*.

Nos últimos anos, os *smartphones* ou telefones “inteligentes” têm figurado no topo das listas dos dispositivos móveis mais vendidos e populares do mercado. Normalmente o que diferencia um *smartphone* de um telefone celular convencional ou um *feature phone* é a presença ou ausência de um sistema operacional “verdadeiro” gerenciando o dispositivo móvel. Ou seja, um *smartphone* é um dispositivo móvel cuja gerência do hardware e do software instalados é feita por um sistema operacional, com todos os subsistemas básicos implementados, como sistema de gerenciamento de memória, sistema de gerenciamento de processos, sistema de gerenciamento de entrada e saída, etc. O fato desses aparelhos serem gerenciados por um sistema operacional os torna muito mais úteis e versáteis, podendo ser incluídos na categoria de computadores “verdadeiros”, de fato.

Outra categoria muito popular atualmente são os *tablets*. Estes se diferenciam dos *smartphones* geralmente por algumas características, como a presença de telas maiores (em geral telas a partir de 7 polegadas) e o fato de que, na maioria dos casos, não oferecem capacidades de uso das redes de telefonia móvel celular, permitindo acesso à internet e às redes locais através especialmente de conexões Wi-Fi e Bluetooth. As telas sensíveis ao toque (*touch screen*) facilitam o uso de muitas aplicações móveis, como os leitores de notícias, os leitores de livros eletrônicos e a execução de jogos eletrônicos.

Com o aumento do tamanho das telas dos *smartphones*, começaram a ser disponibilizados aos usuários dispositivos com telas entre 6 e 7 polegadas, o que permitiu o surgimento de uma nova categoria de aparelhos móveis, conhecida como *phablets*. Isto é, um dispositivo móvel que possui, na verdade, características tanto dos *smartphones* quanto dos *tablets*. De fato, um *phablet* é um dispositivo grande demais para ser chamado de *smartphone* e pequeno demais para ser chamado de *tablet*.

Finalmente, com a crescente miniaturização dos sensores, telas, processadores e memória, um novo tipo de dispositivo móvel está se tornando popular no mercado. Estes aparelhos podem ser acoplados ao próprio corpo de seus usuários sob a forma de relógios, pulseiras, óculos, roupas, calçados e até tatuagens. A tais dispositivos móveis, dá-se o nome de dispositivos vestíveis (*wearable devices*).

4.3. Computação Ubíqua

Visualize o seguinte cenário: você acorda com o som de sua música favorita e, enquanto toma banho, seu café da manhã já está sendo preparado, na cozinha, por sua assistente pessoal. Esta assistente não é um ser humano, mas entende o que você gosta de comer quando acorda. Enquanto faz sua primeira refeição do dia, as principais notícias da manhã, bem como informações sobre a previsão do tempo, o trânsito e seus principais compromissos do dia são exibidas no espelho que fica ao lado da mesa, na cozinha. Por

comando de voz, você pede para falar com sua secretária, no trabalho, informando-a da confirmação de todos os seus compromissos profissionais de hoje.

Em seguida, vai para a sua corrida matinal realizada todas as manhãs no parque próximo à sua casa. Durante sua atividade física cotidiana, uma série de informações sobre você são obtidas e enviadas a um sistema de gerenciamento e análise de informações pessoais que é responsável por analisar e gerenciar dados sobre sua saúde, desempenho em atividades físicas, viagens, compromissos, dentre outras.

Durante o trajeto de casa para o trabalho, seu automóvel autônomo escolhe a melhor rota e o transporta com segurança e conforto até a empresa onde você trabalha, enquanto você continua a leitura das notícias importantes e se diverte com a central multimídia do veículo. Ao chegar à sua sala, seu computador é ligado automaticamente e as notícias que você ainda estava lendo pelo sistema multimídia de seu automóvel são exibidas também na tela do computador em sua mesa de trabalho, permitindo a você finalizar a leitura de seu jornal da manhã.

Quando chega o momento da sua primeira reunião do dia, novamente sua assistente pessoal digital o alerta sobre a pauta e os participantes da reunião. Em seguida, ao chegar à sala de reuniões, seu *smartphone* altera automaticamente o tipo de toque de campanha para o status “silencioso”, evitando qualquer constrangimento, caso alguma chamada seja feita durante essa atividade. Isso foi possível porque uma aplicação de gerenciamento em seu *smartphone* tinha acesso à sua agenda de compromissos e preferências, permitindo automatizar adequadamente essa e muitas outras tarefas cotidianas.

Ao final da tarde, ao chegar em casa novamente, o sistema de climatização de sua residência altera a temperatura de sua casa de modo a ficar de acordo com suas preferências. Pouco tempo depois que chega em casa, você nota a presença de um entregador do supermercado mais próximo trazendo-lhe vários alimentos que acabaram em seu refrigerador e você nem havia notado. Na verdade, a compra foi realizada pelo próprio refrigerador, identificando a ausência de vários alimentos que você consome regularmente e fazendo a compra automaticamente pela internet.

É importante notar que este cenário, para muitos, futurista e distante, já não é possível apenas nos sonhos ou na imaginação visionária de alguns escritores e cineastas. Também não é privilégio de poucos pesquisadores, que trabalham dia e noite em laboratórios de pesquisa de novas tecnologias muito avançados. Pode-se dizer que, com os avanços da Computação Móvel e da Computação Ubíqua, esse mundo tão desejado por Mark Weiser [Weiser 1991] e outros cientistas e pesquisadores já está se tornando realidade em muitos lugares do planeta.

4.3.1 Introdução à Computação Sensível ao Contexto

Entretanto, para que estes e outros cenários semelhantes se tornem realidade, os sistemas computacionais devem se comportar de forma mais autônoma, “inteligente”. Isto só é possível se forem inseridas, nestes sistemas, capacidades que normalmente são encontradas nos seres vivos, algumas delas, quase que exclusivamente nos seres

humanos. Dentre essas capacidades, estão a possibilidade de sentir o ambiente ao seu redor e reagir às mudanças ocorridas nele.

Estes sistemas computacionais com essas características são chamados de *Sistemas* ou *Aplicações Sensíveis ao Contexto*, porque são capazes de perceber as mudanças nas informações de contexto, como o nível de energia do dispositivo, a velocidade e a latência das redes de comunicação, a localização dos seus usuários/dispositivos e se adaptar a essas mudanças de modo a serem mais úteis e autônomos a seus usuários.

O emprego da sensibilidade ao contexto na Computação Móvel permite novas possibilidades no desenvolvimento de sistemas ou aplicações mais adaptadas a esse ambiente ubíquo altamente mutável e inconstante em que estão imersos os dispositivos móveis. Porém, apesar das possibilidades oferecidas, essa nova forma de computação também traz novos e muitos desafios a serem superados. Alguns dos mais comuns são a grande quantidade de tecnologias de sensoramento, como a possibilidade de coletar a localização de um usuário através de vários tipos de sensores e formatos, e também o grande número de informações de contexto (uso de CPU, localização dos dispositivos e usuários, agenda de compromissos, preferências, etc.)

Tudo isso mostra que desenvolver aplicações sensíveis ao contexto sem o emprego de infraestruturas de provisão de contexto (*middleware*, *frameworks*, etc.) adequadas pode ser uma tarefa extremamente complexa, pois pode exigir que o desenvolvedor entenda e programe tanto os subsistemas de coleta, de processamento e o de alteração de comportamento das respectivas aplicações computacionais.

É claro que, se o desenvolvedor pudesse se dedicar exclusivamente às tarefas mais diretamente relacionadas à lógica da aplicação em si, seu trabalho seria mais eficiente, pois não teria de gastar tanto tempo com tarefas relativas à provisão de informações de contexto. Além das questões relativas ao desenvolvimento das aplicações, um outro ponto negativo está relacionado à manutenção de tais sistemas. Por exemplo, se um sensor for substituído por outro equivalente, porém com interface de acesso aos dados diferente, a equipe de desenvolvimento deverá dispender talvez uma quantidade considerável de tempo e esforço para permitir que o novo dispositivo se comunique adequadamente com a aplicação já existente.

Desse modo, as pesquisas, o desenvolvimento e a disponibilização de infraestruturas para a provisão de informações de contexto, como *toolkits*, *middleware* e *frameworks* especializados são vitais para o sucesso e a popularização da Computação Sensível ao Contexto e, por conseguinte, da própria Computação Ubíqua. Isso ocorre porque tais infraestruturas provêm mecanismos que facilitam tanto a coleta, quanto o processamento e a disponibilização de dados de contexto às aplicações, diminuindo o trabalho dos desenvolvedores, aumentando a qualidade dos sistemas desenvolvidos e encorajando novos profissionais a criarem este tipo de aplicação ou sistema.

Como exemplo, há na literatura vários trabalhos que analisam e relatam o desenvolvimento de várias plataformas de provisão de contexto, cada uma com suas particularidades, funcionalidades e finalidades específicas. Dentre estas plataformas ou infraestruturas de apoio ao desenvolvimento de aplicações sensíveis ao contexto, pode-

se citar o Context Toolkit [Salber 1999] [Dey 2001], o ConBus [Sá 2010], a MoCA [Sacramento 2004] [MoCA Team 2005] e o Hydrogen [Hofer 2003].

A definição de Contexto é essencial para o estudo e o desenvolvimento de aplicações e sistemas *sensíveis ao contexto*. Portanto, será apresentada, na subseção seguinte, uma discussão sobre as diversas definições do termo “contexto” nas mais variadas áreas do conhecimento, bem como a evolução dessa definição dentro da Computação Ubíqua.

4.3.2. O que é Contexto

A expressão “*sensibilidade ao contexto*” foi utilizada, pela primeira vez, segundo Dey *et al* [Abowb 1999], em 1994, em um trabalho dos pesquisadores B. Schilit e M. Theimer [Schilit 1994]. Para estes autores, a sensibilidade ao contexto era a capacidade que uma aplicação ou sistema possuía de se adaptar às pessoas e aos objetos co-localizados, bem como à variação de posição dessas pessoas e objetos com o passar do tempo, incluindo também a capacidade desse software identificar e se adequar ao local onde a aplicação estava sendo executada.

Apesar da expressão “*sensibilidade ao contexto*” ter sido usada em textos científicos apenas em 1994, os primeiros sistemas sensíveis ao contexto são anteriores a esta data. Dois anos antes, ainda em 1992, um trabalho desenvolvido no centro de pesquisas da Olivetti e denominado *Active Badge Location System* [Want 1992], é considerado o primeiro sistema baseado em localização conhecido, criado a partir da visão futurista de Mark Weiser e seus companheiros.

Para Seng W. Loke [Loke 2006], o emprego do termo “contexto” tem sido usado em diversas áreas do conhecimento humano, desde a linguística, passando pela teoria da comunicação, até a filosofia e a inteligência artificial. Por isso, pode-se encontrar inúmeras definições dessa palavra, dependendo da área de interesse.

Assim, contexto pode ser “o encadeamento de ideias de um escrito”, conforme o Dicionário Michaelis [Michaelis 2016] ou “aquilo que está em volta, e dá significado a algo”, de acordo com o *Free Online Dictionary of Computing* [Dictionary 2011]. As duas definições anteriores não ajudam muito os pesquisadores de Computação Sensível ao Contexto, pois são, de maneira geral, muito amplas e vagas para serem empregadas nesta área.

Desse modo, a primeira definição mais próxima dos sistemas distribuídos e das redes de computadores é aquela encontrada no trabalho já mencionado acima de M. Theimer e B. Schilit [Schilit 1994], seguida também por uma definição semelhante de Ryan *et al.* [Ryan 1997]. Nestas definições, o termo contexto é definido basicamente como sendo a informação de localização, bem como a identificação das pessoas e objetos co-localizados e suas respectivas mudanças. Esta ideia de contexto já apresenta um grande avanço para a área, porém ainda não é muito operacional e prática para a Computação Sensível ao Contexto e para a Computação Ubíqua em geral.

Uma definição mais adequada de contexto vem através do trabalho de Dey *et al.* [Abowd 1999] e descreve contexto como sendo “*qualquer informação que possa ser usada para caracterizar a situação de entidades (por exemplo, pessoas, locais ou objetos) que são consideradas relevantes para a interação entre um usuário e uma*

aplicação, incluindo também o usuário e a própria aplicação. A Figura 4.3 representa graficamente essa definição.

Pelo fato dessa definição se adequar aos objetivos desse campo do conhecimento, ela será a definição empregada no restante desse capítulo, sempre que for empregado o termo “contexto”.

4.3.3. Classificação de Contexto

A classificação de contexto pode ser feita segundo vários aspectos e critérios. Serão analisados a seguir algumas dessas classificações.

Para T. Gu *et al.* [Gu 2005], há dois tipos principais de contexto: *contexto direto* e *contexto indireto*.

Neste caso, *contexto direto* está relacionado a toda informação de contexto obtida ou adquirida diretamente de um provedor de contexto (serviço especializado a fornecer informações de contexto às aplicações) e pode ser ainda subdividido em contexto que pode ser sentido (*sensed context*) e contexto definido (*defined context*).

No caso de contexto que pode ser sentido (*sensed context*), um provedor de contexto pode fornecer informações coletadas diretamente de *sensores físicos*, incluindo os dados de localização de um usuário através do receptor GPS de seu dispositivo móvel e dos valores de temperatura e luminosidade de uma sala de aula ou ainda pode coletar dados de *sensores virtuais*, como as informações sobre a temperatura de uma região ou cidade distante obtida através de um serviço web específico (*web service*). Quando se trabalha com o contexto definido (*defined context*) se está interessado naquelas informações “definidas” pelo próprio usuário da aplicação, como as preferências de configuração de seu dispositivo móvel, a sua agenda de compromissos ou até mesmo a lista de restaurantes preferidos.

Por outro lado, o *contexto indireto* é sempre obtido através de alguma interpretação de outros contextos diretos ou indiretos (*context reasoning*). Como exemplo, pode-se inferir, com certa margem de erro, através da interpretação de contexto, se uma pessoa está ou não tomando banho verificando-se sua localização (dentro ou fora do banheiro), se a porta está ou não fechada e se o chuveiro está ou não ligado.

A Figura 4.4 ilustra as diferentes categorias de contexto discutidas nos parágrafos anteriores.

Outra forma de classificar contexto pode ser vista no trabalho de Schilit *et al.* [Bill 1994]. Para estes pesquisadores, pode-se categorizar contexto como pertencente a um dos três grupos seguintes: contexto físico, contexto computacional e contexto do usuário. Esta categorização está ilustrada na Figura 4.5.

Contexto Físico diz respeito às grandezas físicas em geral como temperatura, níveis de ruído, luminosidade, dentre muitas outras.

Contexto Computacional está relacionado aos recursos computacionais de cada dispositivo, como a quantidade de memória disponível, a velocidade de transmissão de dados e a latência da rede de comunicação, o uso de CPU, etc.

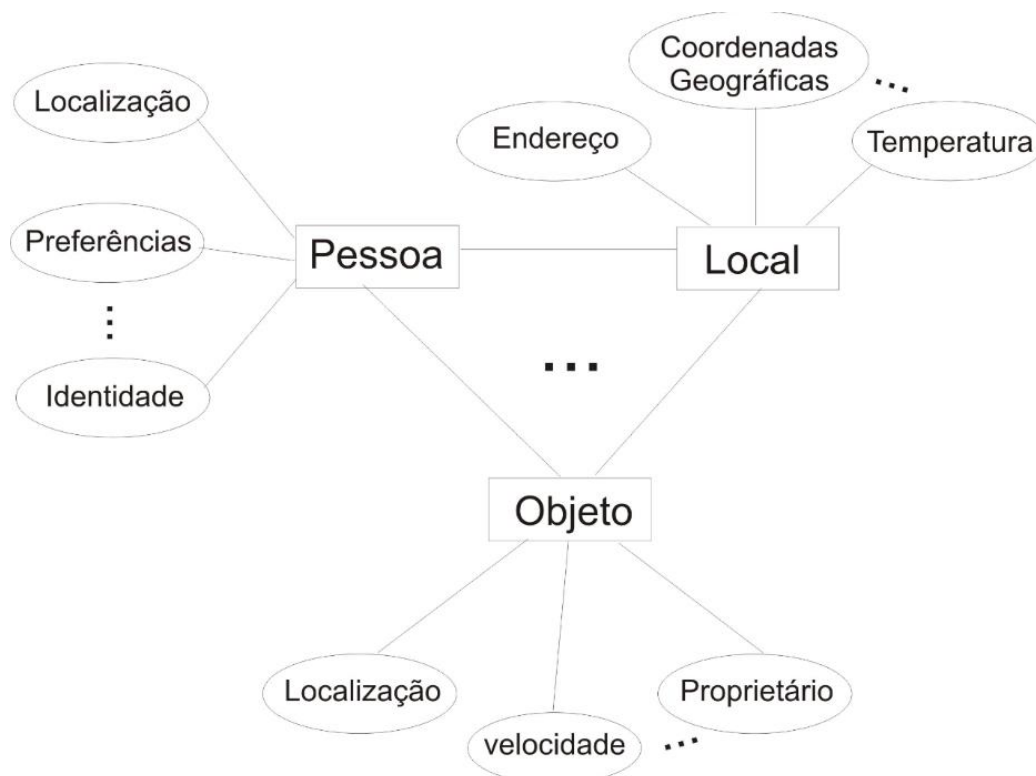


Figura 4.3 - Definindo contexto, segundo [Abowd 1999]

Contexto do Usuário se refere a questões específicas do próprio usuário da aplicação, como sua localização, as pessoas co-localizadas, as preferências e o perfil deste usuário.

Uma outra categoria ainda é proposta por Guanling Chen e David Kotz [Chen 2000]: o *contexto temporal*, ou seja, a estação do ano, o dia da semana, a hora do dia, o mês atual.

4.3.4. Sistemas Baseados em Localização - LBS

Os Sistemas Baseados em Localização (*Location-based Systems* – LBS) são caracterizados pelo uso intensivo e essencial da informação de localização. Neste caso, a informação sobre “onde” o usuário/dispositivo se encontra em cada momento faz toda a diferença para o comportamento da aplicação.

Um dos grandes desafios para os sistemas baseados em localização é garantir a identificação da localização de seus usuários a cada instante, tanto em ambientes externos (*outdoors*) quanto dentro de construções, como casas, escritórios, lojas, shopping centers (*indoors*) [Deak 2012]. Esta preocupação se deve ao fato de que a fonte de informação de localização mais empregada para ambientes externos, o receptor GPS, não é muito adequada para localizar usuários e dispositivos em ambientes internos. Neste caso, se faz necessário lançar mão de outras técnicas e tecnologias mais indicadas para essas situações, como a triangulação de antenas de redes Wi-Fi.

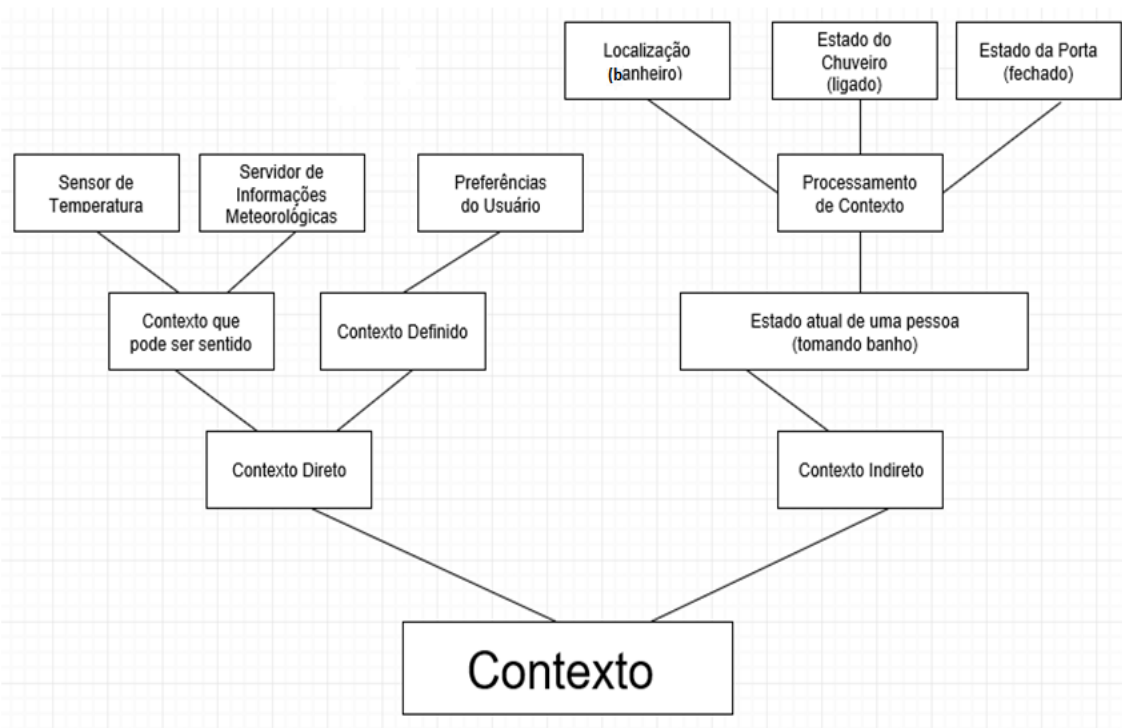


Figura 4.4 – As categorias de contexto definidas por [Gu 2005]

Dentre as aplicações LBS mais populares atualmente, destacam-se os sistemas de navegação empregados especialmente nas grandes cidades, que auxiliam os motoristas a evitarem as vias mais congestionadas ou a encontrar de forma mais eficiente um determinado endereço. Além desses, estão se tornando muito comuns, os sistemas de localização de produtos e serviços próximos aos usuários, como aplicativos de compra de comidas e bebidas, de escolha de táxi, dentre outros.

4.3.5. Computação Vestível

O que se convencionou denominar Computação Vestível (*Wearable Computing*) é, na verdade, uma área tecnológica que objetiva embutir sistemas computacionais sensíveis ao contexto em roupas, relógios, óculos e outros acessórios utilizados normalmente no dia a dia dos seres humanos [Aleksy 2011]. Apesar de estar se popularizando nos últimos anos, podem ser encontradas pesquisas sobre dispositivos vestíveis ainda no século passado [Schmidt 1999].

Estas tecnologias empregam pequenos sensores capazes de coletar as informações contextuais importantes para cada dispositivo, processá-las local ou remotamente, dependendo da complexidade e dos objetivos de cada aplicação. Além disso, são capazes de trocar dados entre dispositivos e aplicações, através de redes sem fio/internet, de modo a criarem redes e sistemas mais complexos que podem realizar diferentes tarefas para seus usuários, como o gerenciamento de informações sobre sua saúde, qualidade de vida, atividades físicas, etc.

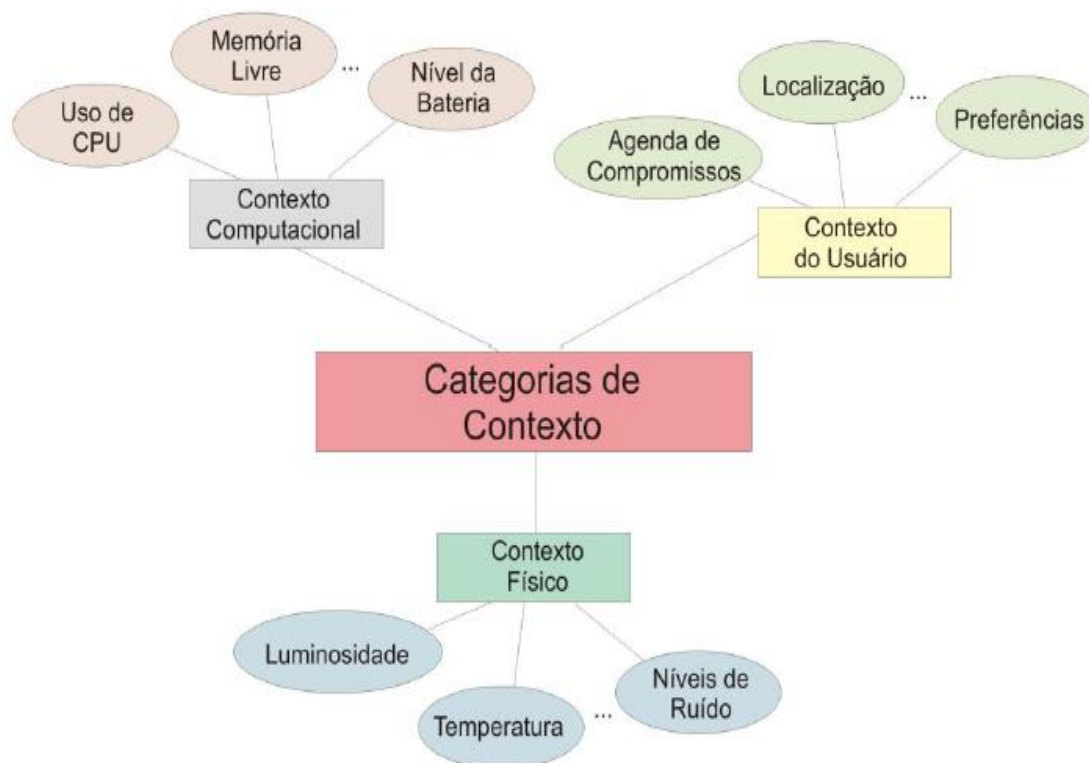


Figura 4.5 – Outra forma de categorizar contexto, definida por [Bill 1994]

4.3.6. A Internet das Coisas

A Internet das Coisas (*Internet of Things* – IoT) é um novo campo de pesquisas e desenvolvimento das novas tecnologias de informação e comunicação (TICs) que visa conectar qualquer “coisa” digital ou física, através da internet, de modo a permitir o desenvolvimento de aplicações e serviços sem precedentes na história da Computação [Razzaque 2016].

Pode-se dizer que a internet passou por três grandes fases ou gerações durante a sua curta história.

A primeira fase, conhecida por *Internet das Máquinas*, diz respeito aos primórdios da internet, quando o objetivo principal ainda era conectar as máquinas, os computadores. Nesta fase praticamente não havia dispositivos móveis.

A partir da década de 1990, com a popularização da web, tem-se o que se chamou de *Internet das Pessoas*, pois o foco muda das máquinas para as pessoas, os usuários, os seres humanos. Neste momento, há o aparecimento de várias redes sociais e o uso da web para realizar atividades financeiras e comerciais. Principalmente a partir de 2010, a internet começou a mudar novamente. Agora, além de interconectar máquinas e pessoas, é potencialmente possível interconectar qualquer “coisa”, ou seja, objetos e dispositivos do dia a dia. Assim, pode-se dizer que o tempo da *Internet das Coisas* está chegando e que, por isso, o mundo tecnológico está passando por um período de transição.

As estimativas são de que em 2020, existam cerca de 50 bilhões de “coisas” conectadas à internet¹. Esta data coincide com as expectativas da entrada em operação das redes de comunicação de Quinta Geração da Telefonia Móvel (5G). Assim, a grande necessidade de comunicação poderá ser suprida pelas novas redes, permitindo o aumento no volume de informações trocadas pelos seres humanos e também por todas as “coisas” conectadas à internet.

No contexto da Internet das Coisas, podem-se definir dois conceitos importantes: “*coisa*” e *dispositivo*, a saber:

- “*coisa*” pode ser qualquer objeto da vida cotidiana, isto é, uma “coisa” pode ser um refrigerador, uma fechadura de porta, um carro, uma residência ou escritório.
- *dispositivo* pode ser um sensor, uma etiqueta ou um atuador. Em geral, um dispositivo é parte integrante de uma “coisa”.

Assim, uma “coisa” é qualquer objeto que, de alguma forma, possui utilidade para os seres humanos, mas que por serem acoplados a algum dispositivo computacional com capacidade tanto de processamento quanto de comunicação, pode interagir e se comunicar com outras “coisas”, através principalmente da internet.

Em relação à arquitetura de sistemas que fazem uso da Internet das Coisas, há, na literatura, várias opções e modelos de referência, como o sugerido por Distefano *et al.* [Distefano 2015], que é exibido na Figura 4.6.

Segundo os autores, a arquitetura de referência para Internet das Coisas (IoT-A) proposta neste trabalho [Distefano 2015] é decomposta em sete grupos de componentes. Cada grupo é discutido em maiores detalhes a seguir:

- *Gerenciamento de Processo de Internet das Coisas*: fornece as funcionalidades relativas à integração dos sistemas de gerenciamento de processos (de negócios) com a infraestrutura de Internet das Coisas.
- *Organização de Serviço*: atua como um “*hub*” de comunicação entre os outros grupos de funcionalidades.
- *Entidade Virtual*: mantém e organiza as informações relacionadas a entidades físicas, habilitando a descoberta de serviços que exibem recursos associados a entidades virtuais e físicas.
- *Serviço de Internet das Coisas*: fornece as funcionalidades requeridas pelos serviços para processamento de informação e para notificar aplicações e serviços sobre eventos relacionados a recursos e entidades físicas correspondentes.
- *Comunicação*: fornece o conjunto de primitivas para a conectividade e comunicação do dispositivo tanto quanto para o roteamento baseado em conteúdo, fornecendo uma interface comum para os serviços de Internet das Coisas.
- *Gerenciamento*: gerencia recursos eficientemente em termos de custo, eventos inesperados, manipulação de erros e flexibilidade.

¹ <http://www.cisco.com/cisco/web/UK/solutions/trends/iot/portfolio.html>

- *Segurança*: é encarregado de assegurar privacidade e segurança aos sistemas que estão em conformidade com a arquitetura IoT-A proposta no respectivo trabalho.

Ainda segundo os autores, os dois últimos grupos (Gerenciamento e Segurança) implementam funcionalidades relacionadas a questões de interesse de todos os outros serviços, e portanto, são verticais, de modo a interfacear com todas as outras camadas.

4.4. Considerações Finais

É evidente que ainda há um longo caminho para que a Computação Ubíqua se torne uma realidade para a grande maioria da humanidade, porém muito já foi feito para isso e com o auxílio das novas tecnologias da informação e comunicação, como a Computação Vestível, os Sistemas Baseados em Localização e a própria Internet das Coisas, o mundo caminha rapidamente para que o computador se torne realmente uma tecnologia invisível, de fato, como vislumbrado por Mark Weiser e outros pesquisadores que também compartilharam deste mesmo sonho ainda nas décadas de 1980 e 1990.

Não há dúvida de que esta nova forma de computação será um dos pilares das tecnologias de informação e comunicação deste século XXI, permitindo o desenvolvimento de inúmeras aplicações que até muito pouco tempo atrás eram inimagináveis para a maioria da população, bem como para a maioria dos próprios programadores e desenvolvedores.

Porém, ainda há muitos desafios para serem superados, como a total implantação do IPv6, a melhoria dos algoritmos e sistemas de segurança e privacidade dos dados, o barateamento dos dispositivos e tecnologias vestíveis e a melhoria da velocidade de transmissão de dados de muitas redes de comunicação espalhadas pelo mundo inteiro.

Com todos esses desafios e problemas, espera-se que novos pesquisadores e profissionais se especializem em Computação Ubíqua, trabalhando em suas subáreas, desenvolvendo aplicações e sistemas sensíveis ao contexto capazes de usar os recursos oferecidos pelos dispositivos vestíveis e a Internet das Coisas.

Desse modo, o computador poderá se transformar em um dispositivo que estará em qualquer lugar a qualquer momento sob as mais variadas formas e tamanho, auxiliando os seres humanos em suas atividades cotidianas pessoais e profissionais, tornando-se assim um dispositivo mais “inteligente” e útil a seus usuários.

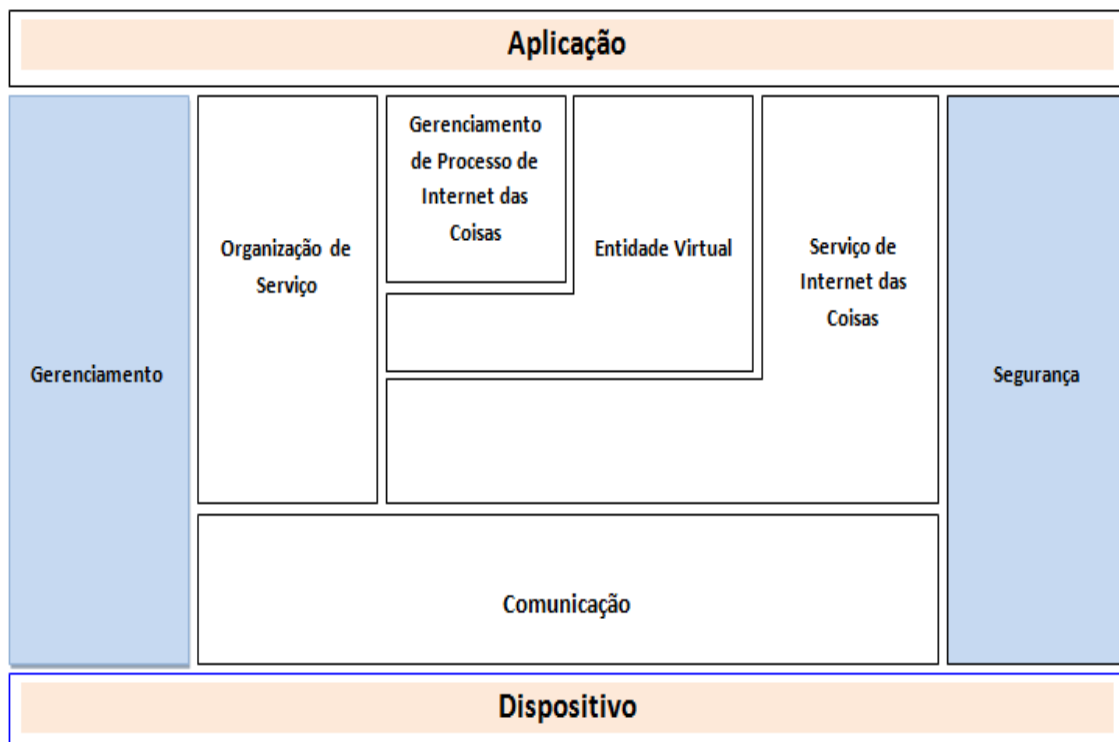


Figura 4.6 – Arquitetura de Referência para Internet das Coisas (simplificada)

Fonte: Traduzida de [Distefano 2015]

Referências

- Razzaque, Mohammad Abdur and Milojevic-Jevric, Marija and Palade, Andrei and Clarke, Siobhán. (2016). “Middleware for Internet of Things: A Survey”. In: IEEE Internet of Things Journal. Volume 3, Issue 1. <http://ieeexplore.ieee.org/xpl/articleDetails.jsp?arnumber=7322178>
- Distefano, Salvatore and Merlino, Giovanni and Puliafito, Antonio. (2015). “A utility paradigm for IoT: The sensing Cloud, Pervasive and Mobile Computing”, Volume 20, July 2015, Pages 127-144, ISSN 1574-1192, <http://dx.doi.org/10.1016/j.pmcj.2014.09.006>.
- Deak, Gabriel. and Curran, Kevin. and Condell, Joan (2012). “A survey of active and passive indoor localisation systems”. In: Computer Communications, 35(16): 1939 – 1954.
- Aleksy, Markus and Rissanen, Mikko J. and Maczey, Sylvia and Dix, Marcel. (2011). “Wearable Computing in Industrial Service Applications”, Procedia Computer Science, Volume 5, Pages 394-400, ISSN 1877-0509, <http://dx.doi.org/10.1016/j.procs.2011.07.051>.
- Schmidt, Albrecht and Beigl, Michael and Gellersen, Hans-W (1999). “There is more to context than location”. Computers & Graphics, Volume 23, Issue 6, Pages 893-901, ISSN 0097-8493, [http://dx.doi.org/10.1016/S0097-8493\(99\)00120-X](http://dx.doi.org/10.1016/S0097-8493(99)00120-X).

- Loureiro, A. A. F. and Mateus, G. R. (1998) “Introdução à Computação Móvel”, In: 11a Escola de Computação, COPPE/Sistemas, NCE/UFRJ., Brazil.
- Weiser, M. (1991) “The computer for the twenty-first century”, In: M., Scientific American, pp. 94–100, September.
- Loke, S. (2006) “Context-Aware Pervasive Systems”, In: Auerbach Publications, Boston, Ma, USA.
- Ryan, J. and Pascoe, N. and Morse, D. (1997) “Enhanced reality fieldwork: the context-aware archaeological assistant,” In: Gaffney, V., Leusen, M. v. and Exxon, S. (Eds.).
- Abowd, G. D. and Dey, A. K. and Brown, P. J. and Davies, N. and M. Smith and P. Steggles (1999) “Towards a better understanding of context and context-awareness,” In: HUC '99: Proceedings of the 1st international symposium on Handheld and Ubiquitous Computing. Springer-Verlag, pp. 304–307, London, UK.
- Schilit, B. N. and Theimer, M. M. (1994) “Disseminating active map information to mobile hosts,” In: IEEE Network, 8(5), pages 22-32. [Online]. Available: <http://schilit.googlepages.com/ams.pdf>.
- Want, R. and Hopper, A. and Falcão, V. and Gibbons, J. (1992) “The active badge location system,” In: Olivetti Research Ltd. (ORL), 24a Trumpington Street, Cambridge CB2 1QA, Tech. Rep. 92.1. [Online]. Available: [iteseer.nj.nec.com/want92active.html](http://citeseer.nj.nec.com/want92active.html).
- Michaelis, Dicionário. (2016) ”Michaelis: Moderno Dicionário de Português Online. In: [Online]. Available: <http://michaelis.uol.com.br/moderno/portugues/>.
- Dictionary. (2011) “Free online dictionary of computing,” In: [Online]. Available: <http://dictionary.reference.com>.
- Gu, T. and Pung, H. K. and Zhang, D. Q. (2005) “A service-oriented middleware for building context-aware services,” In: Journal of Network and Computer Applications, vol. 28, no. 1, pp. 1–18, January. [Online]. Available: <http://dx.doi.org/10.1016/j.jnca.2004.06.002>.
- Bill N. A. and Schilit N. and W. R. (1994) “Context-aware computing applications,” In: Proceedings of the Workshop on Mobile Computing Systems and Applications, Santa Cruz, CA, 8(5), pages 85-90, IEEE Computer Society. [Online]. Available: <http://schilit.googlepages.com/publications>.
- Chen, G. and Kotz, D. (2000) “A survey of context-aware mobile computing research,” In: Technical Report TR2000-381 - Dartmouth College. [Online]. Available: <http://www.cs.dartmouth.edu/reports/TR2000-381.pdf>.
- Salber, D. and Dey, A. K. and Abowd, G. D. (1999) “The context toolkit: aiding the development of context-enabled applications,” In: CHI '99: Proceedings of the SIGCHI conference on Human factors in computing systems. New York, NY, USA: ACM, pp. 434–441.
- Dey, A. and Salber, D. and Abowd, G. (2001) “A conceptual framework and a toolkit for supporting the rapid prototyping of context-aware applications”. [Online]. Available: <http://citeseer.ist.psu.edu/dey01conceptual.html>.

Sá, M. P. de (2010) “Conbus: Uma plataforma de middleware de integração de sensores para o desenvolvimento de aplicações móveis sensíveis ao contexto,” In: Dissertação de mestrado, Instituto de Informática – Universidade Federal de Goiás (INF/UFG), Brasil.

MoCATeam (2016) “Moca home page,” In: <http://www.lac.inf.puc-rio.br/moca> (Last visited: April 2016).

Sacramento, V. and Endler, M. and Rubinsztein, H. K. and Lima, L. S. and Gonçalves, K. and Nascimento, F. N. and Bueno G. A. (2004) “Moca: A middleware for developing collaborative applications for mobile users” In: IEEE Distributed Systems Online, vol. 5, no. 10, p. 2,

Promoção



Organização



Cooperação



Patrocínio



Apoio

