



INSTITUTO
UNIVERSITÁRIO
DE LISBOA

Manutenção Prescritiva: Uso de PLN para identificação e prevenção de incidentes em pagamentos

João Mário Deitado Saraiva

Mestrado em Business Analytics

Orientadores:

Doutor Ricardo Daniel Santos Faro Marques Ribeiro, Professor Associado, ISCTE - Instituto Universitário de Lisboa

Doutor Eugénio Alves Ribeiro, Professor Auxiliar, ISCTE - Instituto Universitário de Lisboa

Julho, 2026

Departamento de Métodos Quantitativos para Gestão e Economia

Manutenção Prescritiva: Uso de PLN para identificação e prevenção de Incidentes em pagamentos

Mestrado em Business Analytics

Orientadores:

Doutor Ricardo Daniel Santos Faro Marques Ribeiro, Professor Associado, ISCTE - Instituto Universitário de Lisboa

Doutor Eugénio Alves Ribeiro, Professor Auxiliar, ISCTE - Instituto Universitário de Lisboa

Julho, 2026

Dedicatória

Ao meu grande avô, Mário Deitado, por todo o amor e carinho que sempre demonstraste por mim, pela felicidade e alegria que sempre proporcionaste a todos à tua volta, pelas brincadeiras e gargalhadas que partilhei contigo e que irei guardar para sempre com muito carinho.

Esta página foi intencionalmente deixada em branco

Agradecimentos

A presente dissertação é o termo de uma longa jornada, que não teria sido possível sem o contributo de muitas pessoas.

Antes de mais, gostaria de agradecer aos meus orientadores, Professor Ricardo Ribeiro e Professor Eugénio Ribeiro, pela orientação constante, pela ajuda incansável e pela sua disponibilidade e paciência ao longo de todo este percurso.

Um agradecimento especial ao Professor Raúl Laureano, pela paciência e pelo tempo dedicado ao apoio prestado durante a realização deste trabalho.

A todos os docentes com quem tive oportunidade de aprender no ISCTE, e que proporcionaram os alicerces de conhecimento essenciais à concretização deste estudo.

Aos meus pais e irmãos pela presença constante, pelo apoio incondicional, pela compreensão, pela paciência e pelos bons momentos partilhados, que foram essenciais para recarregar as energias durante o tempo que durou este projeto académico.

Esta página foi intencionalmente deixada em branco

Resumo

Os incidentes operacionais afetam as instituições financeiras, estando frequentemente associados a impactos na reputação das organizações envolvidas. Estes incidentes operacionais resultam em atrasos na realização de pagamentos ou até em pagamentos para a entidade errada, podendo levar a multas, processos, perda de clientes ou mesmo dos fundos transferidos. As organizações financeiras baseiam a sua atividade na confiança dos clientes. Nesse sentido, a boa reputação é fundamental para o funcionamento e sustentabilidade do negócio. Por esse motivo, os incidentes operacionais têm um impacto substancial nas organizações e são um problema sensível para as organizações financeiras.

De forma a reduzir o impacto negativo dos incidentes operacionais numa organização financeira, há a necessidade de aumentar a quantidade de recursos de remediação dos mesmos, de maneira a assegurar a prevenção de incidentes. Para tal, este projeto propõe complementar o processo manual de classificação de incidentes com um processo automatizado, através de técnicas de Processamento de Linguagem e descoberta de conhecimento em texto, com o objetivo de melhorar o entendimento sobre as causas responsáveis por incidentes na organização.

Para desenvolver este projeto analítico foi utilizada a metodologia *Cross Industry Standard Process for Data Mining* (CRISP-DM), sendo esta a metodologia mais comum para projetos de *Data Mining*. Esta metodologia divide o projeto em seis fases sequenciais: compreensão do negócio, compreensão dos dados, preparação dos dados, modelação, avaliação e implementação, permitindo ao investigador alternar entre estas conforme as necessidades e desafios do projeto.

No contexto deste projeto, os incidentes analisados ocorrem em equipas operacionais responsáveis pelo processamento e manutenção de instruções de pagamento em sistemas financeiros, estando frequentemente associados a erros de configuração, falhas na interpretação de instruções ou atrasos na atualização de informação crítica. Os incidentes considerados no presente estudo são auto-reportados pela equipa em formato de relatório, contendo predominantemente dados não-estruturados sob a forma de descrições textuais das causas dos incidentes. Foram utilizadas ferramentas de *Text Mining* para a extração e classificação de padrões a partir do texto, recorrendo a modelos de *Topic Modeling* probabilísticos e baseados em *embeddings*. Por fim, foi realizado um inquérito de validação dos resultados ao universo operacional da equipa de produção, com o objetivo de mitigar possíveis vieses gerados pela análise de modelos não supervisionados.

O modelo final identificou 18 tópicos relacionados com as causas dos incidentes. As equipas de produção avaliaram os resultados como claros e relevantes para a discussão e prevenção de risco dentro da organização, vendo relevância na expansão do estudo com dados anuais.

Palavras-chave: Incidentes operacionais; *Settlements*; Instituição Financeira; Risco operacional; *Text Mining*; *Topic Modelling*; CRISP-DM; *Business Analytics*; Dados não-estruturados.

Classificação JEL: G00; D53

Esta página foi intencionalmente deixada em branco

Abstract

Operational incidents affect financial institutions and are often associated with reputational impacts on the organisations involved. These incidents may result in delays in payment execution or even in payments being sent to the wrong beneficiary, potentially leading to fines, legal proceedings, loss of clients, or even the loss of transferred funds. Financial institutions operate on the basis of trust, which is essential for clients to entrust them with their money. In this context, a strong reputation is fundamental to the functioning and sustainability of the business. For this reason, operational incidents have a substantial impact on organisations and represent a particularly sensitive issue for financial institutions.

In order to reduce the negative impact of operational incidents within a financial organisation, it is necessary to strengthen remediation resources so as to ensure more effective incident prevention. To this end, this project proposes complementing the manual incident classification process with an automated approach based on Natural Language Processing and Text Mining techniques, with the aim of improving the understanding of the underlying causes of incidents within the organisation.

The analytical development of this project followed the Cross-Industry Standard Process for Data Mining (CRISP-DM), one of the most widely adopted methodologies for Data Mining projects. This methodology structures the project into six sequential phases: business understanding, data understanding, data preparation, modelling, evaluation, and deployment, allowing the researcher to iterate between phases according to the needs and challenges encountered throughout the project.

In the context of this project, the incidents analysed occur within operational teams responsible for the processing and maintenance of payment instructions in financial systems, and are often associated with configuration errors, misinterpretation of instructions, or delays in updating critical information. The incidents analysed in this study are self-reported by the operational team in the form of incident reports, consisting predominantly of unstructured data in the form of textual descriptions of incident causes. Text Mining techniques were applied to extract and classify patterns from the textual data, using both probabilistic topic modelling approaches and embedding-based models. Finally, a validation survey was conducted among the operational production teams to assess the results and mitigate potential biases inherent to the analysis of unsupervised models.

The final model identified 18 topics related to the causes of operational incidents. The production teams evaluated the results as clear and relevant for the discussion and prevention of operational risk within the organisation, recognising the value of extending the analysis to future years using updated data.

Keywords: Operational incidents; Settlements; Financial institution; Operational risk; Text mining; Topic modelling; CRISP-DM; Business analytics; Unstructured data.

JEL Classification: G00; D53.

Esta página foi intencionalmente deixada em branco

Lista de Abreviaturas, Acrónimos e Siglas

- **ABE** – Autoridade Bancária Europeia
- **BERT** – *Bidirectional Encoder Representations from Transformers*
- **BPC** – Banco Popular da China
- **C-TF-IDF** – *classed Term Frequency–Inverse Document Frequency*
- **CRISP-DM** – *Cross-Industry Standard Process for Data Mining*
- **DM** – *Data Mining*
- **ESG** – *Environmental, Social and Governance*
- **FXMM** – *Foreign Exchange Money Markets*
- **GLDA** – *Guided Latent Dirichlet Allocation*
- **IA** – Inteligência Artificial
- **IRD** – *Interest Rate Derivatives*
- **ITSM** – *Information Technology Service Management*
- **KG** – *Knowledge Graph*
- **KNN** – *K-Nearest Neighbors*
- **LDA** – *Latent Dirichlet Allocation*
- **LLM** – *Large Language Model*
- **LTCM** – *Long-Term Capital Management*
- **ML** – *Machine Learning*
- **NB** – *Naive Bayes*
- **NLP** – *Natural Language Processing*
- **SVM** – *Support Vector Machine*
- **SWIFT** – *Society for Worldwide Interbank Financial Telecommunication*
- **TF-IDF** – *Term Frequency–Inverse Document Frequency*
- **TM** – *Text Mining*
- **TPP** – *Third-Party Payment*
- **TSR** – *Total Share Revenue*
- **YAKE** – *Yet Another Keyword Extractor*
- **PLN** – *Processamento de Linguagem Natural*

Esta página foi intencionalmente deixada em branco

Índice

1	Introdução	1
1.1	Motivação	2
1.2	Problema e questão de investigação	2
1.3	Objetivos e contributos da investigação	2
1.4	Pertinência da investigação	3
1.5	Estrutura da dissertação.....	3
2	Revisão da literatura	5
2.1	Conceitos relevantes.....	6
2.1.1	Definição de risco.....	6
2.1.2	Risco operacional.....	6
2.1.3	Tipos de risco operacional	7
2.1.4	Incidentes operacionais.....	8
2.1.5	<i>Text Mining</i> - definição e aplicações	8
2.2	Trabalho relacionado.....	9
2.2.1	Análise de incidentes operacionais em sistemas de pagamentos.....	9
2.2.2	Análise de incidentes operacionais	9
2.2.3	Aplicações de <i>Topic Modeling</i> na análise de incidentes organizacionais e financeiros	10
2.3	Considerações finais	14
3	Compreensão dos dados	15
3.1	Recolha dos dados	15
3.2	Variáveis disponíveis	15
3.3	Tratamento dos dados para a modelação.....	16
3.3.1	Tratamento de nulos.....	17
3.3.2	Variáveis removidas.....	17
3.4	Tratamento de valores omissos e inválidos.....	17
3.5	Variáveis adicionadas.....	18
3.6	Tratamento de dados não estruturados para a geração de tópicos	18
3.7	Caracterização dos dados textuais	20
3.8	Análise exploratória dos dados	20
3.9	Considerações finais	27
4	Modelação	29
4.1	<i>Latent Dirichlet Allocation</i>	29
4.1.1	Abordagem com 3 tópicos	31
4.1.2	Abordagem com 6 tópicos	33

4.1.3	Abordagem com 8 tópicos	34
4.1.4	Abordagem com 10 tópicos	35
4.1.5	Considerações	36
4.2	<i>BERTopic</i>	37
4.2.1	<i>BERTopic</i> utilizando <i>HDBSCAN</i>	38
4.2.2	Análise complementar com apoio de LLM	47
4.2.3	Considerações	49
5	Validação dos Resultados através de Inquérito às Equipas.....	51
5.1	Descrição do questionário e fundamentação das questões	51
5.2	Análise e interpretação das respostas	52
5.3	Considerações finais	56
6	Conclusões	57
6.1	Contribuições	57
6.2	Limitações do projeto	57
6.3	Melhorias futuras	58
7	Referências Bibliográficas.....	61
8	Anexo	65

Índice de Tabelas

Tabela 1: Frequência de valores omissos por variáveis	17
Tabela 2: Normalização de termos e expressões no processo de tratamento textual	19
Tabela 3: Distribuição das 10 moedas com maior número de incidentes	20
Tabela 4: Distribuição de incidentes por causa	23
Tabela 5: Distribuição do número de incidentes por causa e ocorrência de pagamento por engano	24
Tabela 6: Distribuição dos incidentes por mês em frequência e valor monetário	26
Tabela 7: Avaliação de métricas extrínsecas entre os modelos BERTopic	41
Tabela 8: Frequência e Proporção de Documentos por Tópico	42
Tabela 9: Avaliação de resultados feita pelo GPT5.1	48

Esta página foi intencionalmente deixada em branco

Índice de Figuras

Figura 1: Distribuição do número de palavras por documento	20
Figura 2: Distribuição da frequência e do valor dos Incidentes por software	21
Figura 3: Distribuição proporcional do valor em EUR por causa de incidente, segmentado pelo sistema responsável	22
Figura 4: Distribuição e impacto financeiro dos incidentes por causa	23
Figura 5: Distribuição proporcional dos montantes pagos ("Amount in EUR") por causa do incidente ("Cause") e ocorrência de pagamento indevido ("Paid by Mistake")	24
Figura 6: Distribuição proporcional do montante em EUR por causa do incidente e por responsabilidade	25
Figura 7: Distribuição de incidentes por dias da semana e valor dos incidentes em euros por dia da semana	26
Figura 8: Variação da coerência em função do número de tópicos no modelo LDA	30
Figura 9: Nuvens de palavras representativas dos três tópicos gerados pelo modelo LDA	31
Figura 10: Redes de conceitos associadas aos três tópicos gerados pelo modelo.	31
Figura 11: Mapa de Componentes Principais e termos mais relevantes	32
Figura 12: Distribuição percentual de documentos por tópico no modelo LDA com três tópicos	33
Figura 13: Termos mais relevantes por tópico no modelo BERTopic com 3 tópicos	38
Figura 14: Estrutura hierárquica dos clusters de tópicos gerados pelo modelo BERTopic com 3 tópicos	39
Figura 15: Mapa de Similaridade entre tópicos no modelo BERTopic com 3 tópicos	40
Figura 16: Termos mais relevantes por tópicos no modelo BERTopic com 18 tópicos	43
Figura 17: Mapa de Calor da Similaridade entre Tópicos	45
Figura 18: Dendrograma de Agrupamento Hierárquico dos Tópicos	45
Figura 19: Projeção Bidimensional dos Documentos e Distribuição dos Tópicos	47
Figura 20: Avaliação da adequação do título atribuído a cada tópico face ao conteúdo dos incidentes	53
Figura 21: Avaliação da clareza dos tópicos com base nos termos dominantes e incidentes representativos	54
Figura 22: Identificação de causas de risco reveladas pelo modelo "Análise por conjuntos de resposta e por tópicos"	55
Figura 23: Avaliação global da qualidade da análise - Classificação média atribuída pelos inquiridos	56

Esta página foi intencionalmente deixada em branco

1 Introdução

Os incidentes operacionais nos sistemas de pagamento podem ter impactos significativos na estabilidade financeira e no funcionamento dos mercados (Russo e Steconci, 2011).

Sistemas de pagamento são uma parte vital da economia e das instituições financeiras, dado que o seu funcionamento eficiente permite que as transações sejam processadas atempadamente e de forma segura, contribuindo para o desenvolvimento da economia (Bank of England, 2000).

Com o crescimento do volume de pagamentos, a adoção de tecnologias mais sofisticadas e a transição dos meios tradicionais em papel para sistemas eletrônicos, o acesso facilitado da população a esses serviços têm redefinido o risco operacional. Esse cenário tem tornado os bancos mais vulneráveis a incidentes operacionais severos (Heinrich, 2006).

Os incidentes de produção financeira representam uma preocupação significativa para instituições financeiras (Heinrich, 2006), refletindo-se em perdas substanciais que podem afetar a sua estabilidade e reputação. Incidentes operacionais em pagamentos podem acontecer por variados motivos, tais como erro humano, erro do sistema ou falhas em processos operacionais (SWIFT, 2025).

Continua a ser um desafio classificar de forma correta as principais causas que levam à ineficiência dos processos de pagamentos, seja pela multiplicidade das possíveis causas que levam ao incidente, seja pelas regras extensas e complexas dos *softwares* de produção ou mesmo pelas próprias regras complexas da rede *Society for Worldwide Interbank Financial Telecommunication* (SWIFT).

A consciência sobre o risco operacional tem aumentado ao longo dos anos, tanto no nível individual quanto no âmbito das instituições financeiras, abrangendo os sistemas de pagamento, compensação e liquidação (Heinrich, 2006).

Essas preocupações intensificaram-se com a crescente automação e integração dos sistemas. Um incidente isolado em um sistema automatizado, pode desencadear efeitos negativos significativos em diversas instituições, resultando em uma cadeia complexa de falhas nos processos de pagamento (Miller & Northcott, 2002).

Entre 2016 e 2021, incidentes operacionais, tiveram um impacto na indústria financeira em cerca de 600 mil milhões de euros, sendo este um aumento de quase três vezes comparado ao período de 2011 e 2016 (ORX, 2022). Só em 2023 foram perdidos 15.2 mil milhões de euros em eventos de risco operacional no setor financeiro (ORX, 2024). Ainda assim, estes números estão consideravelmente abaixo da média de 21.5 mil milhões de euros da indústria bancária nos últimos 6 anos. Nos últimos anos, a principal causa de eventos de falhas operacionais foi a fraude externa, com 166.803 ocorrências entre 2018 e 2023 (ORX, 2024).

De acordo com um estudo feito pela McKinsey (2023) a 500 eventos de risco operacional, numa fase inicial estes acontecimentos tiveram um impacto no preço das ações em linhas com as multas, os pagamentos e as perdas monetárias. No entanto, a longo prazo estas empresas tendem a perder até 2,7% de *Total Shareholder Revenues* (TSR) nos retornos totais até 120 dias após o evento, comparado com empresas semelhantes.

Em 2008, a *Société Générale* enfrentou um dos maiores casos de fraude da história bancária, resultando em um impacto financeiro de quase 5 mil milhões de euros, decorrente de um evento significativo de risco operacional (Banco de Portugal, 2014).

Em junho de 2024, a UBS anunciou que reservou cerca de 840 milhões de euros para compensar investidores afetados pelo colapso da *Greensill Capital*. Esta medida visa reembolsar antigos clientes do *Credit Suisse* que investem em fundos associados à *Greensill*, oferecendo-lhes 90% do valor dos seus ativos perdidos num evento de risco operacional da *Credit Suisse*.

1.1 Motivação

A crescente complexidade dos sistemas financeiros e o aumento do volume de transações processadas diariamente têm reforçado a necessidade de as organizações financeiras desenvolverem mecanismos mais eficazes de monitorização, compreensão e prevenção do risco operacional. Neste contexto, a identificação das principais causas associadas aos incidentes de produção assume um papel particularmente relevante, não apenas pela necessidade de mitigar perdas financeiras e reputacionais, mas também pela importância de garantir a continuidade, eficiência e fiabilidade dos sistemas de pagamento.

Apesar da elevada quantidade de informação produzida diariamente pelas equipas operacionais, uma parte significativa desse conhecimento encontra-se dispersa em dados textuais não estruturados, dificultando a identificação sistemática de padrões recorrentes e das causas latentes dos incidentes. Assim, surge a necessidade de desenvolver abordagens que permitam transformar estes dados em informação útil para apoio à decisão e melhoria contínua dos processos operacionais.

Neste contexto, o presente projeto procura responder a uma necessidade organizacional concreta, através da extração automática de tópicos e padrões recorrentes, pretende-se contribuir para o aumento do conhecimento sobre as causas dos incidentes, facilitando a discussão de medidas preventivas, a priorização de ações corretivas e a identificação de áreas críticas de risco dentro da organização.

1.2 Problema e questão de investigação

A principal questão de investigação formulada para o presente projeto foi a seguinte:

Como melhorar a compreensão sobre os incidentes de produção? Mais especificamente: como é que a aplicação de técnicas de *Text Mining* (TM) em dados textuais pode melhorar a identificação de padrões e a compreensão das causas dos incidentes operacionais em sistemas de pagamentos?

A resposta a esta questão, bem como à subquestão dela derivada, procura demonstrar de que forma a aplicação de ferramentas de TM aos relatórios de incidentes pode contribuir para a identificação de tópicos comuns entre as falhas operacionais, permitindo às equipas elaborar planos de ação orientados para a prevenção da recorrência desses incidentes, com base na informação extraída dos tópicos

1.3 Objetivos e contributos da investigação

Este estudo tem como principais objetivos a extração de tópicos relevantes a partir da documentação associada a incidentes, recorrendo a técnicas de TM, com vista à identificação das causas latentes mais recorrentes.

Pretende-se, adicionalmente, avaliar a eficácia e a aplicabilidade desta abordagem em contexto real, através de um inquérito de validação de resultados dirigido aos colaboradores das equipas onde os dados foram recolhidos, com o objetivo de verificar de que forma os resultados da metodologia respondem à questão de investigação e ao problema de negócio. Por fim, ambiciona-se a definição de um processo sistemático de extração de informação a partir de dados não estruturados, que possa ser reaplicado e ajustado em análises futuras, à medida que novos dados forem sendo disponibilizados.

1.4 Pertinência da investigação

Este estudo visa, de forma geral, contribuir para o enriquecimento da literatura sobre incidentes operacionais em sistemas de pagamentos, através da aplicação de ferramentas de TM para identificar as principais causas latentes associadas às falhas operacionais. Verifica-se, atualmente, a existência de uma lacuna na literatura científica no que diz respeito à análise sistematizada de incidentes operacionais neste domínio específico, bem como à utilização de técnicas de TM aplicadas a este tipo de dados. Assim, esta investigação pretende colmatar esse vazio, fornecendo um contributo original e metodologicamente fundamentado.

Ficará, assim, disponível para investigadores e estudos futuros uma proposta de metodologia que contemple os vários aspetos do processamento de linguagem natural, bem como um modelo cuja finalidade seja a identificação de tópicos relevantes no contexto operacional dos sistemas de pagamentos.

Para além disso, espera-se que sejam obtidas conclusões orientadas para a liderança do departamento das equipas analisadas, permitindo-lhes conhecer os tópicos associados aos incidentes financeiros, bem como o seu impacto em termos de frequência e valor monetário. Desta forma, a liderança poderá identificar as principais fontes de risco no seio das equipas e tomar decisões informadas que previnam a recorrência desses incidentes, ficando a saber à partida quais as causas a priorizar no caso de ser necessário desenvolver projetos de melhoria de *software*, os quais implicam a alocação de recursos temporais e financeiros.

1.5 Estrutura da dissertação

Este trabalho está estruturado em seis capítulos principais, seguindo a metodologia *cross-industry standard process for data mining* (CRISP-DM), amplamente utilizado em projetos de análise de dados, que estrutura o processo em seis fases: compreensão do negócio, compreensão dos dados, preparação dos dados, modelação, avaliação e implementação. No contexto deste projeto, estas fases foram aplicadas de forma adaptada. Na fase de compreensão do negócio, foi identificado o desafio de interpretar incidentes operacionais descritos em dados textuais não estruturados. Seguiu-se a compreensão e exploração dos dados disponíveis, analisando a sua estrutura e qualidade. Na fase de preparação, foram realizadas tarefas de pré-processamento textual, incluindo limpeza, normalização e transformação dos dados. Posteriormente, na fase de modelação, foram aplicadas técnicas de topic modeling, tanto probabilísticas como baseadas em embeddings. A fase de avaliação incluiu a análise de métricas de coerência e a validação dos resultados através de um inquérito às equipas operacionais. Por fim, a abordagem foi pensada numa perspetiva de aplicação contínua, permitindo a sua reutilização em análises futuras.

No que diz respeito à estrutura do documento, o primeiro capítulo, correspondente à introdução, apresenta a contextualização do tema, a definição do problema de investigação e a justificação da sua relevância.

O segundo capítulo é dedicado à revisão da literatura, onde são enquadrados os principais conceitos relevantes e as abordagens metodológicas associadas à análise de incidentes

operacionais, com especial enfoque na aplicação de técnicas de text mining. No terceiro capítulo, procede-se à compreensão dos dados, descrevendo o conjunto de dados utilizado e explorando em detalhe as variáveis disponíveis.

O quarto capítulo aborda a fase de modelação, onde são apresentadas e comparadas as metodologias adotadas, nomeadamente o Latent Dirichlet Allocation (LDA) e o BERTopic, bem como os respetivos processos de parametrização e avaliação. O quinto capítulo centra-se na validação dos resultados através de um inquérito às equipas envolvidas nos incidentes de produção, descrevendo o questionário utilizado, a fundamentação das questões e a análise e interpretação das respostas obtidas. Por fim, o sexto capítulo apresenta as conclusões do estudo, incluindo a discussão dos resultados alcançados, as principais limitações identificadas ao longo do projeto e a proposta de possíveis melhorias e desenvolvimentos futuros que permitam aprofundar ou expandir esta investigação.

2 Revisão da literatura

Ao longo deste capítulo serão abordados temas e fontes diferentes, permitindo enquadrar o projeto desenvolvido na instituição financeira, tendo como objetivo classificar os incidentes operacionais em pagamentos, através de ferramentas de TM, permitindo a adoção de medidas preventivas, seja por meio da melhoria de processos, seja pelo desenvolvimento de *softwares* de produção que reduzam a ocorrência desses incidentes.

Dessa forma, este capítulo começará por introduzir conceitos relevantes fundamentais para o entendimento deste projeto, começando por definir o conceito de risco, para de seguida alargar a compreensão sobre o conceito de risco operacional e incidentes operacionais em sistemas de pagamentos. De seguida, serão explorados os principais conceitos relacionados com a metodologia de TM aplicada neste estudo. Por fim, serão discutidas diversas abordagens para a identificação de tópicos nos incidentes, utilizando metodologias de *text clustering*.

A presente revisão de literatura teve como objetivo mapear os principais conceitos, métodos e aplicações de técnicas de *text mining* em incidentes operacionais e contextos financeiros. A pesquisa bibliográfica foi realizada entre janeiro e março de 2025, utilizando bases de dados científicas como Scopus e WoS. Foram utilizadas combinações de palavras-chave relacionadas com *text mining*, *topic modeling*, análise de incidentes e risco operacional, recorrendo a expressões de pesquisa como TITLE-ABS-KEY(("text mining" OR "topic modeling" OR "NLP") AND ("operational incidents" OR "operational risk") AND ("financial systems" OR banking OR payments)). Para garantir uma abrangência adequada, não foram definidos limites cronológicos iniciais, mas priorizou-se a literatura produzida nos últimos 10 anos (2015 a 2025).

Os critérios de inclusão abrangeram artigos revistos por pares, estudos de caso aplicados e investigações empíricas que envolvessem a utilização de *natural language processing* (NLP) em contextos operacionais ou financeiros. Foram também considerados estudos fora do setor financeiro como saúde, construção civil e gestão de incidentes técnicos sempre que as abordagens metodológicas fossem transferíveis para o domínio em análise. A triagem inicial baseou-se na leitura de título e resumo, seguida da leitura integral dos artigos mais relevantes para o presente estudo.

Durante esta pesquisa, verificou-se a existência de uma lacuna significativa no estado da arte relativamente à investigação de incidentes operacionais em sistemas de pagamento, decorrente da escassez de artigos científicos que proponham ou analisem modelos de prevenção destes incidentes. A maioria dos trabalhos identificados aborda o tema sobretudo numa perspetiva jurídica e burocrática, centrando-se na interpretação da legislação que regula a atividade financeira. É possível considerar que esta escassez de publicações decorra da elevada sensibilidade do tema e do impacto que a divulgação de incidentes possa ter na reputação das organizações financeiras envolvidas. Esta circunstância impossibilitou a realização de uma revisão sistemática da literatura, uma vez que não se encontraram estudos suficientemente relevantes que permitissem estabelecer comparações e análises mais aprofundadas. Optou-se, assim, por uma revisão narrativa, direcionada para a análise de incidentes operacionais noutros setores, cujas abordagens podem ser extrapoladas para a realidade dos dados disponíveis e para a questão de investigação em estudo.

2.1 Conceitos relevantes

Para melhor compreender a problemática em análise e sustentar a aplicação de técnicas de TM na classificação de incidentes operacionais em sistemas de pagamento, é necessário, em primeiro lugar, clarificar um conjunto de conceitos essenciais. Este subcapítulo introduz, de forma progressiva, as noções de *risco*, *risco operacional* e *incidentes operacionais*, estabelecendo as bases teóricas que enquadram o presente estudo. A seguir, são apresentados os principais fundamentos associados à metodologia de *text mining*, que constituem a base para a modelação e análise subsequente dos dados.

2.1.1 Definição de risco

O conceito de risco acompanha a humanidade desde os primórdios da Humanidade (Banco de Portugal, 2014). Desde as agressões de predadores até à incerteza de encontrar alimentos, os seres humanos sempre enfrentaram desafios que os obrigaram a avaliar e a gerir o risco das suas ações (Ruppenthal, 2013).

Este conceito é amplamente estudado em diversas disciplinas, desde a economia e gestão até às ciências sociais, devido à sua importância em diferentes contextos. Todos os seres humanos estão expostos ao risco e criam, desenvolvem e acumulam níveis de experiência na sua gestão, com base nos eventos vivenciados ao longo da vida pessoal de cada um (Adams, 2002).

Em 1983, o grupo de estudo *Britain's Royal Society* publicou um relatório chamado de *Risk Assessment*, que se tornou um trabalho de referência na área do risco, que via o risco como a probabilidade de um acontecimento adverso acontecer num determinado período, ou como resultado de um desafio específico. Este relatório faz a distinção entre *risco objetivo* e *risco percebido*, sendo o primeiro uma medição baseada em fatos e análises de modo a identificar e quantificar o risco de forma objetiva e o segundo ao modo como os indivíduos interpretam o risco de forma qualitativa.

De acordo com o Professor e geógrafo John Adams (2002), existe uma associação entre o risco e a incerteza, sendo o risco o conhecimento da probabilidade de algo acontecer, apesar do desconhecimento do que de facto irá acontecer. Para este Professor, o risco é determinado pelo conhecimento das probabilidades de ocorrência de cada acontecimento, caso contrário, estaríamos, segundo ele, a lidar com incerteza.

Para o Professor Eduardo Sá Silva (2015), o risco pode ser definido como a possibilidade de perda, ou o grau de incerteza relacionado com um acontecimento. O Professor considera que é importante diferenciar o risco de incerteza, na medida em que primeiro permite a estimativa objetiva das probabilidades dos acontecimentos, enquanto que, no caso da incerteza, não são possíveis tais estimativas, tendo-se de recorrer a probabilidades subjetivas.

Seguindo o mesmo raciocínio, para alguns autores (Guagliard *et al.*, 2016), o conceito de risco está associado a uma incerteza delineada pela probabilidade mensurável de determinadas ações produzirem danos, afetando os resultados previstos. Para estes autores, existe ainda a distinção entre *risco confirmado* e entre *risco potencial*, sendo o risco confirmado o que é previsível e o potencial aquele que não é capaz de se identificar.

2.1.2 Risco operacional

O risco operacional é uma área em constante evolução, que tem ganhado maior destaque devido ao fortalecimento da cultura de auto-regulação no setor financeiro (Hemrit & Arab,

2013) e a uma série de eventos que evidenciaram a necessidade de maior controlo sobre os processos operacionais, como o colapso do *Baring Banks* em 1995 e a crise do *Long-Term Capital Management* (LTCM) em 1998, tendo a investigação subsequente atribuído as respetivas falhas à gestão inadequada do risco operacional (Peter, Gordon & Yueran, 2018).

Em 2006, o comité de Supervisão Bancária de Basileia emitiu um relatório que define o *risco operacional* como o risco ou a perda de pessoas e sistemas ou de eventos externos resultante de processos internos inadequados ou falhados. Esta definição inclui o *risco legal*, mas exclui o *risco estratégico* e *reputacional* (Basel Committee on Banking Supervision, 2006). Esta abordagem destaca a natureza abrangente do risco organizacional, que pode surgir tanto de falhas nos processos, como erros nos sistemas, erro humano e até eventos externos inesperados. A separação do risco estratégico e reputacional é importante na delimitação no âmbito desta definição.

A Diretiva Europeia 2009/138/CE segue uma variante da definição dada pelo Comité de Basileia na Supervisão Financeira, definindo o *risco operacional* como “risco de perdas resultantes de procedimentos internos inadequados ou deficientes, do pessoal ou dos sistemas, ou ainda de acontecimentos externos”

O Banco de Portugal (2014) define o *risco operacional* como o risco de perdas ou impactos negativos financeiros, no negócio e/ou na imagem/reputação da organização, causados por falhas ou deficiências na governação e processos de negócio, nas pessoas, nos sistemas ou resultantes de eventos externos, que poderão ser despoletados por uma multiplicidade de eventos.

As organizações financeiras estão bastante expostas a risco organizacional dado a sua dependência de atividades e processos sujeitos a erro humano, a falhas de sistemas e de infraestruturas (Jorion, 2007).

No relatório *ORX Banking Operational Risk Loss Data Summary Report 2024*, as principais causas dos eventos de perdas operacionais no setor bancário estão categorizadas como: Fraude Externa; Execução, Entrega e Gestão de Processos; Clientes, Produtos e Práticas de Negócios; Fraude Interna; Falhas de Sistemas; Eventos externos; Proteção de Dados e Segurança e Roubo.

2.1.3 Tipos de risco operacional

Marc C. Dobler *et al.* (2021) distinguem cinco tipos de risco operacional: *fraude interna e externa*, *ciber risco*, *risco de agência*, *risco de negócio* e *risco de investimento*. Já Gutierrez e Jeffrey (2006) consideram ainda o *risco de informação* em adição aos vários tipos de risco operacional.

Ivarinen *et al.* (2003) notam que, no contexto dos sistemas de pagamento, os principais riscos operacionais estão relacionados à informação, tecnologia, administração e criminalidade. Por sua vez, Merrouche e Schanz (2010) analisaram o impacto do risco operacional sobre outros tipos de risco nesses sistemas, com ênfase no risco de liquidez.

O comité de Supervisão Bancária de Basileia define como risco operacional, os riscos de perdas causados por falhas dos processos, falhas de pessoas, de sistemas ou de eventos externos.

2.1.4 Incidentes operacionais

Os sistemas de pagamento são críticos para a estabilidade financeira, e as falhas inerentes podem causar efeitos sistêmicos. Segundo a Autoridade Bancária Europeia (ABE), um *incidente operacional* pode ser definido como um evento singular ou uma série de eventos interligados, não planeados pelo agente do serviço de pagamentos que irão ter um impacto adverso na integridade, disponibilidade, confidencialidade e autenticidade de serviços relacionados com pagamentos (2017).

No contexto dos sistemas de pagamento, o risco operacional é visto como o risco de incorrer em encargos com juros ou com outros custos resultantes de desvios ou falhas na execução atempada de pagamentos no Forex Market Settlement, devido a erro humano ou falha técnica (CPSS, 1996).

2.1.5 *Text Mining* - definição e aplicações

A crescente disponibilidade de grandes volumes de dados não estruturados tem impulsionado o uso de ferramentas de *text mining* em diversos setores, como saúde, finanças, *marketing* e gestão operacional (Khan *et al.*, 2020). Estas ferramentas tornaram-se fundamentais para a extração automatizada de conhecimento e para a interpretação de grandes volumes de texto.

Text Mining pode ser definido como a extração de *insights* a partir de dados não estruturados não só por meio de aprendizagem automática (Aggarwal, 2018), mas também pode ser entendido como o processo de descoberta de padrões e conhecimento em dados não estruturados (Akilan, 2015).

Envolve o processo de compilar, organizar e analisar em conjunto de grandes coleções de documentos para apoiar a entrega de tipos de informação direcionadas a analistas e decisores bem como para encontrar relações entre fatos relacionados que abranjam amplos domínios de investigação (Al-Ghuribi & Noah, 2022).

Embora semelhante ao *Data Mining* (DM), o *Text Mining* diferencia-se pelo tipo de fontes de dados utilizadas, que incluem informações semi-estruturadas ou não estruturadas, como *e-mails*, *tweets*, documentos de texto e ficheiros HTML (Fan, Wallace, Rich, & Zhang, 2005). As suas aplicações abrangem a descoberta de novos padrões em grandes volumes de dados, extração de informação, *clustering* e classificação de documentos, *Web Mining*, extração de conceitos, frequentemente suportadas por técnicas de Processamento de Linguagem Natural (PLN), que permitem representar e transformar texto em variáveis analisáveis (Hotho, Nürnberger, & Paaß, 2005).

Estas ferramentas abrangem um conjunto diverso de abordagens, que podem ser adaptadas a diferentes tipos de dados e objetivos analíticos. Khan *et al.* (2020) classificam essas abordagens em várias categorias principais: (1) extração de informação, focada na identificação automática de entidades e relações nos documentos; (2) classificação de texto supervisionada, usada para categorizar documentos com base em rótulos pré-definidos; (3) *clustering* não supervisionado, que permite agrupar documentos semelhantes entre si, sem conhecimento prévio das categorias; (4) modelação de tópicos e rastreamento de eventos, útil para identificar e acompanhar temas emergentes; e (5) sumarização automática, para a redução de conteúdo textual mantendo a informação essencial.

Apesar de desafios técnicos como a ambiguidade linguística, a alta dimensionalidade e a escassez de contexto em textos curtos, estas técnicas oferecem um leque robusto de

possibilidades analíticas (Khan *et al.*, 2020). Combinadas, permitem transformar grandes volumes de dados textuais em estruturas interpretáveis e úteis, especialmente relevantes em domínios como a análise de incidentes, onde os padrões nem sempre são imediatamente visíveis por métodos tradicionais.

2.2 Trabalho relacionado

Esta secção apresenta as principais abordagens metodológicas para a análise e prevenção de risco relativo a incidentes operacionais identificadas na literatura, organizadas de forma progressiva. Num primeiro momento, são discutidos os artigos especificamente relacionados com incidentes em sistemas de pagamento, ainda que em número reduzido. Segue-se a análise de estudos que abordam incidentes operacionais em outros setores de atividade, cuja transferibilidade metodológica pode oferecer contributos relevantes para o presente projeto. Por fim, são exploradas as diferentes metodologias de *topic modeling* aplicadas à análise de incidentes operacionais, com destaque para os modelos mais recorrentes na literatura recente.

2.2.1 Análise de incidentes operacionais em sistemas de pagamentos

Diversos estudos têm se debruçado sobre a análise de incidentes em sistemas de pagamento, devido ao seu impacto direto na confiança dos consumidores e na estabilidade do sistema financeiro (Heinrich G., 2006).

Por exemplo, Yinhong Yao e Jianping Li (2022) realizaram uma análise quantitativa sobre o impacto dos riscos operacionais em plataformas de *Third-Party Payments* (TPP), utilizando uma abordagem de *Loss Distribution Approach*. Este estudo identificou duas principais fontes de risco operacional em sistemas de pagamento: *Risco Legal*, correspondendo a 95,05% dos eventos de risco e o *Risco de Fraude Externa*, 4,95% dos eventos de risco. O risco legal decorre de multas regulatórias e sanções aplicadas por órgãos como o Banco Popular da China (BPC), por falhas e incumprimentos nas normas financeiras de segurança. A principal causa de risco legal identificado é a falha na identificação de clientes com 26,54% dos eventos, seguida da não submissão de relatórios relativos a transações suspeitas (14,20%).

Por sua vez, Wenhao Kang e Chi Fai Cheung (2023) utilizaram uma abordagem inovadora na análise e identificação de incidentes em sistemas bancários, utilizando *reconhecimento de entidades nomeadas* para o desenvolvimento de *grafos de conhecimento*, utilizando o modelo aperfeiçoado BERT-BiLSTM-CRF para a extração de entidades nomeadas em incidentes bancários com 96,75% de precisão. Neste artigo, o uso de *grafos de conhecimento* permitiu uma análise mais aprofundada das entidades nomeadas, identificando padrões comuns entre diferentes incidentes.

2.2.2 Análise de incidentes operacionais

Apesar de a literatura sobre sistemas de pagamento abordar frequentemente o tema do risco operacional, a maioria da literatura concentra-se em áreas como a deteção e prevenção de fraude, bem como na análise do risco de crédito que, apesar de relevantes, não se enquadram necessariamente no tema da presente investigação. Para além disso, muitos artigos adotam uma abordagem voltada para o cumprimento regulatório formal e para a gestão burocrática e jurídica de incidentes, analisando a sua conformidade com normas e processos internos, em vez de buscar uma compreensão aprofundada dos padrões e causas subjacentes aos incidentes operacionais.

Mesmo quando são analisados dados de incidentes em sistemas de pagamento, a metodologia tende a basear-se em classificações manuais ou em dados estruturados previamente categorizados. Dados não estruturados, como relatórios ou registros de ocorrências, são pouco explorados.

Neste sentido, existe uma lacuna na literatura relativa à aplicação de técnicas de PLN na análise de incidentes em sistemas de pagamentos, incidentes estes que poderiam proporcionar *insights* valiosos sobre os padrões dos incidentes. De modo a suprir esta lacuna na literatura, o presente projeto procura contribuir para superar essa lacuna, propondo o uso de abordagens de TM que já foram utilizadas para analisar e identificar incidentes operacionais em diferentes setores.

No setor da construção civil, Goh e Ubeynarayana (2017) aplicaram técnicas de TM para automatizar a classificação de narrativas de acidentes, analisando 1.000 relatos de incidentes de construção. O estudo comparou seis classificadores, destacando o SVM linear com o método de *tokenização unigram* e *term frequency-inverse document frequency* (TF-IDF) como o mais eficaz, com *F1-scores* superiores a 0,80 em categorias como “eletrocuções” e “quedas”.

Para melhorar a categorização de incidentes no contexto de *IT Services Management* (ITSM), Sara Silva (2018) utilizou ferramentas de PLN, TM e *Machine Learning* (ML), substituindo o processo manual anteriormente realizado. Para isso, utilizou algoritmos como *Support Vector Machines* (SVM), *K-Nearest Neighbors* (KNN), *Naïve Bayes* (NB) e Árvores de Decisão, obtendo uma precisão de 89%.

Também no contexto de IT, Jorge Costa (2019) desenvolveu um modelo de automatização da resolução de incidentes através de diversos algoritmos de ML como SVM, KNN, NB, Regressão Logística, e Redes Neurais, com o objetivo de prever a categoria de resolução dos *tickets* com base nas descrições textuais, obtendo uma precisão de 58,63%, o que representa uma economia significativa de tempo e recursos para a organização.

Em 2021, Liu *et al.* utilizaram ferramentas de TM para desenvolver um modelo de classificação de incidentes de risco ferroviários, utilizando Bi-LSTM-CRF e *random forest* para extrair e classificar as entidades, de modo a desenvolver um grafo de conhecimento capaz de identificar as principais características dos incidentes operacionais ferroviários, contribuindo a medidas preventivas.

Pedro *et al.* (2022) desenvolveram um sistema de partilha de informação baseado em ontologias e grafo de conhecimento para melhorar a comunicação e prevenção dos incidentes de segurança na construção civil, superando a limitação tradicional de falta de padronização dos dados.

2.2.3 Aplicações de *Topic Modeling* na análise de incidentes organizacionais e financeiros

A modelação de tópicos tem vindo a destacar-se como uma ferramenta promissora na análise automatizada de grandes quantidades de dados textuais, particularmente em contextos onde os incidentes organizacionais ou operacionais são descritos de forma narrativa e não estruturada.

Este subcapítulo procura explorar o uso desta técnica em dois domínios interligados: a análise de incidentes organizacionais em setores diversos e a sua aplicação mais específica no setor financeiro. Através da análise de estudos recentes que empregam algoritmos como o LDA e

o BERTopic, e variantes baseadas em *transformers*, *embeddings* ou ambos, pretende-se perceber se existe evidência suficiente quanto à sua eficácia, relevância e transferibilidade para o contexto dos incidentes operacionais em sistemas de pagamento.

Esta abordagem justifica-se pela necessidade de compreender melhor o potencial destas metodologias e apoiar, com base na literatura existente, a escolha dos modelos mais adequados para a investigação presente.

2.2.3.1 Abordagens baseadas em LDA e variantes

Entre estes algoritmos, o LDA destaca-se como um método não supervisionados de modelação de tópicos que identifica padrões em documentos sem qualquer classificação prévia, atribuindo palavras a tópicos de forma probabilística e aleatória (Blei, Ng e Jordan 2003). Em estudos aplicados a incidentes e a grandes volumes de texto não estruturado, surgem também variantes do LDA que procuram aumentar a relevância e a coerência dos tópicos em domínios específicos.

No setor financeiro, o LDA é igualmente aplicado em corpora extensos, com especial incidência em notícias e sustentabilidade. Kim *et al.* (2024) analisaram mais de 140.000 notícias de 2019 a 2022 relacionadas com a sustentabilidade financeira, com o objetivo de identificar os temas principais ligados à *environmental, social, and governance* (ESG). Os autores utilizaram o modelo LDA com a seguinte *pipeline* de pré-processamento:

1. Remoção de *stopwords* e pontuação;
2. *Tokenização*;
3. *Lematização*
4. Conversão para minúsculas.

Os tópicos extraídos dos títulos incluíram áreas como gestão económica e regulamentar, imagem pública e tecnologia verde.

Yang *et al.* (2022) aplicaram uma metodologia de *text mining* para identificar tópicos e sentimentos associados ao tema da propriedade intelectual na China, nas redes sociais Weibo e WeChat. Os autores utilizaram uma combinação de TF-IDF, *TextRank* e LDA para extração e modelação de tópicos, e uma rede BiLSTM com incidência na análise de sentimentos, obtendo 16 tópicos principais relacionados com o tema. O *Term Frequency–Inverse Document Frequency* (TF-IDF) foi utilizado para identificar as palavras com maior relevância estatística com base na sua frequência relativa nos documentos, enquanto o *TextRank* avalia a importância semântica com base nas relações de co-ocorrência entre os termos, permitindo a seleção das mais representativas e relevantes de palavras-chave, que foram usadas pelo modelo probabilístico de geração de tópicos LDA.

Numa lógica semelhante de combinação de técnicas de extração de palavras-chave com modelação probabilística de tópicos, Ahadh *et al.* (2021) utilizaram ferramentas de *text mining* para analisar relatórios incidentes no setor da aviação e no setor de transportes de materiais perigosos, utilizando o algoritmo *Yet Another Keyword Extraction* (YAKE) para identificar os termos relevantes dos documentos e o *Guided Latent Dirichlet Allocation* (GLDA) para a formação de tópicos, obtendo uma precisão de 80%. O *Guided LDA* surge como um método aperfeiçoado do LDA, na medida em que utiliza palavras-chave predefinidas, garantindo que os tópicos gerados tenham maior relevância e coerência com o setor em análise.

Para além do LDA, alguns trabalhos no domínio financeiro comparam e estendem abordagens probabilísticas com modelos que incorporam *embeddings*, Giaconia *et al.* (2024) exploraram o uso de técnicas de modelação de tópicos para a análise de revisões humanas de alertas de lavagem de dinheiro, no setor da auditoria bancária, aplicando o modelo probabilístico LDA e dois modelos de *embeddings* distintos, o ETM e ProdLDA a 35 relatórios de alertas de branqueamento de capitais. Os modelos foram avaliados utilizando o *score* de coerência, o *score* de NPMI e a diversidade de tópicos. Os autores concluíram que o modelo ProdLDA com *embeddings* Word2Vec gram gerou os tópicos mais coerentes e relevantes para o contexto financeiro. Os tópicos extraídos refletiram atividades específicas, como transações internacionais para países de risco, jogos de azar, relações familiares e pagamento em dinheiro, indicando o potencial da modelação de tópicos para identificar padrões de risco financeiro.

De forma complementar, Masson e Paroubek (2024) investigaram a eficácia de diferentes abordagens de modelação de tópicos na identificação de riscos a partir de relatórios anuais da indústria financeira francesa, utilizando tanto modelos probabilísticos como LDA e NMF, como também modelos de *embeddings* como o ProdLDA e o *Scholar*. O modelo *Scholar*, com covariáveis (c-Scholar) demonstrou o melhor desempenho geral, com maior diversidade, menor viés setorial e maior capacidade de capturar riscos específicos e comuns. Os resultados foram incorporados numa ferramenta de apoio à análise regulatória.

Finalmente, surgem abordagens híbridas que procuram combinar representações semânticas com a extração de distribuições de tópicos do LDA. Para analisar os tópicos latentes em notícias sobre finanças, Zhou *et al.* (2022) propuseram o modelo BERT-LDA, uma abordagem híbrida de modelagem de tópicos que combina *embeddings* semânticos gerados com MacBERT (Cui *et al.*, 2020) e RoBERTa (Liu *et al.*, 2019) com distribuições de tópicos extraídas por LDA. Utilizaram o *classed Term Frequency–Inverse Document Frequency* (c-TF-IDF) para representar semanticamente os tópicos, demonstrando um valor de coerência semântica superior ao LDA e NMF, destacando-se o modelo MacBERT-LDA. Apesar dos resultados positivos dos modelos de *embeddings* contextuais, a representação dos tópicos feita pelo c-TF-IDF é representada por *bag-of-words*, evidenciando a redundância nas palavras que caracterizam o tópico.

2.2.3.2 Abordagens baseadas em BERTopic e modelos com *embeddings* contextuais

Uma das principais limitações do LDA na modelação de tópicos é não considerar a relação semântica entre as palavras. Por utilizar uma representação de *bag-of-words*, os tópicos são gerados de forma probabilística com base na distância entre as palavras, ignorando o contexto das palavras numa frase, podendo levar a uma má representação dos documentos (Grootendorst, 2022). Em resposta a esta limitação, modelos baseados em *embeddings* têm ganho popularidade na modelação de tópicos, como o *Bidirectional Encoder Representations from Transformers* (BERT) (Devlin *et al.*, 2018), mostrando resultados muito relevantes a gerar tópicos bem representados contextualmente (Grootendorst, 2022). Entre estes, destaca-se o BERTopic que combina a representação de *embeddings* que permitem identificar coerência semântica com uma adaptação do TF-IDF chamada de *class-based* TF-IDF (Grootendorst, 2022).

No setor bancário, Ogunleye *et al.* (2023) compararam diferentes abordagens de modelação de tópicos aplicadas ao setor bancário, analisando 10.000 *tweets* de clientes de bancos nigerianos, comparando modelos tradicionais como LDA, LSI e HDP, com variantes do

BERTopic. O BERTopic demonstrou um *score* de coerência de 0,8463, muito superior aos dos modelos tradicionais que variavam entre 0,3919 e 0,6347.

A utilização do BERTopic estende-se também à análise de incidentes em setores operacionais. Com o objetivo de identificar padrões e preocupações públicas relativas à segurança ferroviária, Ghosh *et al.* (2024) analisaram 93.239 *tweets* de janeiro de 2017 a maio de 2022, relacionados com incidentes ferroviários. Encontraram 5 tópicos principais dentro do conjunto de dados analisados através do algoritmo *BERTopic*, sendo estes cultura de segurança, ensino da segurança ferroviária às crianças, uso de *headphones*, caminhar em caminhos ferroviários interditos e notificações de emergência.

Em contexto clínico-hospitalar, os autores Denecke e Paula (2024) utilizaram técnicas de PLN, incluindo o BERTopic para analisar 10.000 relatos de incidentes críticos em hospitais suíços de 2006 a 2023, identificando os principais erros cometidos por médicos, enfermeiro ou máquinas e equipamento hospitalares. Através da modelação dos tópicos concluiu-se que as principais causas estavam relacionadas com medicação como, por exemplo, erros de prescrição, administração de insulina, heparina e cálcio, além de problemas registados em procedimentos laboratoriais, emergências pediátricas, quedas de pacientes, anestesia e incidentes associados à COVID-19. Ainda no setor clínico aplicaram o BERTopic à análise de relatórios clínicos de 2.389 pacientes atendidos por médicos de família nos Países Baixos, entre 2016 e 2019, com o objetivo de identificar padrões preditivos de quedas em idosos. Através do modelo de *embeddings* clínicos MedRoBERTa.nl, foram identificados 25 tópicos que revelavam sinais antecipados de risco como o uso de medicamentos, hospitalizações e necessidade de cuidadores.

Num enquadramento mais abrangente e temporal, Mishra *et al.* (2024) realizaram uma análise temporal extensiva da literatura em economia computacional, aplicando e comparando os modelos LSA, LDA e BERTopic na identificação de tópicos emergentes e padrões de evolução. Utilizando uma base de dados composta por quase 3.000 resumos de artigos científicos publicados de 2000 a 2022, os autores concluíram que o modelo BERTopic demonstrou ser mais confiável na coerência, relevância e distinção temática, validando assim o seu uso em contextos complexos e multidisciplinares. A *pipeline* do BERTopic incluiu a configuração padrão do modelo utilizando *embeddings* gerados com SBERT, redução de dimensionalidade com UMAP, agrupamento com HDBSCAN e extração de palavras-chave com c-TF-IDF, sem necessidade de *fine-tuning* adicional. Entre os temas em crescimento destacam-se: teoria dos jogos, design de mecanismos dinâmicos e precificação de ativos, ilustrando a capacidade da modelação de tópicos em captar tendências de investigação ao longo do tempo.

2.2.3.3 Abordagens com LLMs e GenAI na triagem e interpretação de incidentes

Recentemente *large language models* (LLMs), como o GPT-3.5 e GPT-4, têm demonstrado ser alternativas promissoras às abordagens tradicionais de *text mining*, oferecendo uma maior capacidade de interpretação (Cui *et al.*, 2023). Em contexto operacional, diversos estudos têm demonstrado o potencial destas ferramentas para melhorar a triagem de incidentes (Wang *et al.*, 2024), conseguindo até fornecer soluções orientadas para a prevenção destes (Smetana *et al.*, 2024).

Na análise e identificação de incidentes na indústria de construção rodoviária Smetana *et al.* (2024) utilizaram modelos LLM como o GPT-3.5 para resumir os conteúdo dos *clusters*, identificar as causas mais comuns dentro de cada grupo e classificar ou reavaliar as

categorias existentes. Neste estudo, os tópicos são detetados por um modelo de *embeddings text-embedding-ada-002*, da OpenAI, onde a LLM interpreta e rotula os tópicos.

De forma semelhante, Wang *et al.* (2024) propõem um sistema que utiliza modelos de LLM para realizar a triagem de incidentes técnicos em ambientes de computação em *cloud* na *Microsoft Azure* de forma automática, processando os dados técnicos e os *tickets* de incidentes para identificar os responsáveis e sugerir ações de prevenção. Este sistema é composto por dois módulos principais, o *AutoExtractor*, responsável por identificar palavras-chave e tópicos relevantes nos dados e o GPT-3.5, que interpreta essas informações para apoiar a decisão de triagem. Os testes demonstraram que este sistema reduziu o tempo de resposta em 35% e aumentou a precisão da classificação de incidentes em 30%.

Por sua vez, Otal *et al.* (2024) utilizaram LLMs como o LLaMA2 e *Mistral* para extrair informações-chave no processo de triagem de chamadas de emergência na área da gestão de incidentes e emergências, recolhendo informação crítica em tempo real. Este modelo foi integrado numa aplicação que fornece orientação personalizada com base nas suas mensagens e localização.

2.3 Considerações finais

A literatura analisada evidencia o crescente interesse na aplicação de técnicas de *text mining* para compreender padrões complexos em dados não estruturados, com particular destaque para a modelação de tópicos para identificar informação em grandes volumes de dados não estruturados. Embora existam aplicações bem-sucedidas em áreas como Saúde, Transporte e Serviços de IT, são ainda escassos os estudos que explorem essas técnicas no contexto específico dos sistemas de pagamento e da análise de incidentes operacionais financeiros.

Identificou-se uma clara lacuna no uso de abordagens como o BERTopic e modelos baseados em *embeddings* para analisar relatos de incidentes em serviços financeiros. Nesse sentido, a presente dissertação procura contribuir para colmatar essa lacuna, propondo a aplicação de técnicas avançadas de *topic modeling* à análise de incidentes operacionais em sistemas de pagamento, com vista a melhorar a categorização e compreensão dos mesmos, utilizando modelos probabilísticos, modelos de *embeddings* e *large language models*.

3 Compreensão dos dados

Após a definição do problema de investigação e do enquadramento teórico associado ao presente projeto, torna-se também necessário compreender os dados utilizados no desenvolvimento deste. Esta etapa assume um papel fundamental, permitindo identificar a estrutura, qualidade e relevância da informação disponível, bem como enquadrar este projeto no contexto da organização.

Neste capítulo são apresentados os dados recolhidos, as variáveis disponíveis e os procedimentos adotados para a sua preparação e transformação. Adicionalmente, é realizada uma análise exploratória com o objetivo de identificar padrões iniciais e criar suposições que poderão ser testados após a modelação.

3.1 Recolha dos dados

Para este projeto de investigação foram analisados os dados agregados dos incidentes relativos ao ano de 2023, reportados por equipas operacionais responsáveis pela inserção e atualização em sistema das instruções de liquidação dos clientes para a compra e venda de produtos financeiros. Os dados foram reportados pelas respetivas equipas quando o incidente foi registado na plataforma e atribuído a uma destas equipas.

Cada incidente é caracterizado por diversas variáveis categóricas e temporais que ajudam a compreender, em detalhe, a natureza do incidente, tais como a data em que o incidente foi reportado e a moeda em que a transação foi efetuada. No entanto, a variável não estruturada "*comment*" acaba por ser a mais relevante para analisar os incidentes financeiros, na medida em que contém uma descrição detalhada e cronológica dos acontecimentos que levaram ao incidente, bem como das suas causas.

Quando um acidente é atribuído a uma das equipas referidas, os colaboradores envolvidos analisam detalhadamente o ocorrido e confirmam se a responsabilidade é aceite ou não pela equipa, justificando a decisão através de um relatório. A partir deste documento, extrai-se a informação sobre o incidente para integrar as bases de dados analisadas no presente projeto.

Ao todo, foram registados 495 incidentes durante o ano de 2023.

3.2 Variáveis disponíveis

A base de dados utilizada neste estudo é composta por 18 variáveis que descrevem os incidentes registados, nomeadamente a variável "*Name*" identifica o *software* onde o incidente ocorreu, enquanto "REF" corresponde ao identificador atribuído pela plataforma responsável pelo seu registo. A variável "*PAID BY MISTAKE*" é binária e indica se o incidente resultou num pagamento efetuado por engano. Já a variável "*Amount in EUR*" representa o valor do incidente expresso em euros, independentemente da moeda original da transação, garantindo uniformidade nos valores registados. A variável "*RESPONSABILITY*", também binária, assinala se a responsabilidade pelo incidente foi atribuída à equipa.

As variáveis temporais registam três momentos distintos: "*DATE*" refere-se à data em que o incidente foi reportado, "*Date of the insertion of the SSI*" indica a data em que as instruções do cliente foram inseridas em sistema e "*VALUE DATE*" corresponde ao dia previsto para a execução do pagamento.

A variável "*Cause*" é categórica e identifica, de forma genérica, a razão subjacente ao erro. Entre as categorias disponíveis, destaca-se "*Late SSI action/Cut off*", que se refere a ações

executadas após o prazo limite para o processamento do pagamento. enquanto "*Notifications/Broadcast*" compreende dois tipos distintos de comunicação:

1. As notificações, isto é, mensagens automatizadas que atualizam as instruções dos fundos de clientes que utilizam a plataforma ALERT, uma solução global para o registo e manutenção de dados de contas de fundos acessível aos bancos com os quais mantém relações;
2. Os *broadcasts*, mensagens emitidas pelos bancos a comunicar alterações nas suas contas ou agentes para determinadas moedas.

Estas comunicações têm impacto em todos os *setups* que utilizem as instruções do banco em questão, exigindo a respetiva verificação e atualização por parte das equipas responsáveis.

A categoria "*Others Responsibility*" agrupa os incidentes cuja responsabilidade é atribuída a terceiros, como o cliente, o *back office* ou o *front office*. A categoria "*Wrong Setup*" abrange situações em que as instruções estavam em desconformidade com as normas da SWIFT ou com os requisitos do *software* utilizado. Por ser uma categoria abrangente, é especialmente importante analisar os seus principais fatores. Por fim, a categoria "*Others*" agrega incidentes que não se enquadram em nenhuma das classificações anteriores.

A variável "*CLIENT*" representa um código interno utilizado pelo banco para identificar os clientes. Estes códigos podem ter 7 ou 12 dígitos: os primeiros sete identificam a entidade principal, e os cinco adicionais referem-se à subentidade ou departamento específico.

A variável "*Comment*" contém uma descrição textual detalhada do incidente escrita em inglês, incluindo os registos das ações envolvidas e as respetivas datas. Por esse motivo, trata-se de uma das variáveis mais relevantes do conjunto de dados, pois fornece informação contextual fundamental sobre o incidente, desde o processo envolvido até à possível origem do problema.

A variável "*CCY*" refere-se à moeda em que a transação foi efetuada.

As variáveis "*CREATOR*" e "*VALIDATOR*" identificam os colaboradores responsáveis pela criação e validação do *setup*, respetivamente. A variável "*AURORA/360*" refere-se à plataforma utilizada pela equipa para o registo dos incidentes, enquanto "*TL that check the incident*" identifica o *Team Leader* que procedeu à respetiva análise.

Por fim, a variável "*CASH OUT*" (Sim/Não) indica a ocorrência de um pagamento efetuado para uma conta errada, sendo este um tipo específico de pagamento indevido. Embora semelhante à variável "*PAID BY MISTAKE*", a principal diferença reside no facto de a primeira implicar obrigatoriamente que o montante tenha sido transferido para uma conta incorreta, enquanto a segunda variável se pode referir a um pagamento incorreto, mesmo que direcionado para a conta certa. Neste contexto, um incidente classificado como "*CASH OUT*" representa o cenário mais crítico, dado o seu potencial para gerar perdas significativas e custos pesados para o banco.

3.3 Tratamento dos dados para a modelação

O processo de preparação dos dados para a modelação constitui uma etapa crítica para o desenvolvimento de modelos analíticos. Envolve a normalização de formatos, a conversão de variáveis categóricas, a criação de variáveis derivadas e a seleção dos atributos mais relevantes. Esta fase visa garantir a consistência, a integridade e a adequação estatística dos dados, de modo a maximizar a capacidade explicativa dos modelos e a evitar reveses durante a aprendizagem.

3.3.1 Tratamento de nulos

A gestão de valores em falta é essencial para preservar a validade dos resultados obtidos durante a modelação. A presença de nulos pode distorcer inferências estatísticas e comprometer o desempenho dos algoritmos. Por este motivo, foi realizada uma análise criteriosa da distribuição de valores ausentes por variável, seguida da aplicação de estratégias específicas de imputação ou exclusão, consoante a natureza dos dados e o impacto estimado na qualidade da informação.

Tabela 1: Frequência de valores omissos por variáveis

Variável	Número de Valores Omissos
Name	0
DATE	0
REF	0
PAID BY MISTAKE	1
Amount in EUR	0
VALUE DATE	0
Date of the insertion of the SSI	7
Ccy	5
CLIENT	1
RESPONSABILITY	0
CAUSE	0
Comment	2
CREATOR	222
VALIDATOR	268
AURORA/360	477
TL that check the incident	285
Comments	482
CASH OUT (Yes/No)	493

Para tratar dos valores omissos presentes na Tabela 1 foram seguidas diferentes estratégias conforme a natureza dos dados em questão. Algumas variáveis são removidas do *dataset* pois não agregam informação relevante para a modelação dos dados.

3.3.2 Variáveis removidas

As variáveis "AURORA/360" e "Ref" foram desconsideradas da análise estatística por se tratar de variáveis de identificação, que não contribuiriam para análises futuras. Decidiu-se também remover a variável "CASH OUT (Yes/No)" da análise, não só por apresentar 493 valores omissos.

Por fim, a variável "DATE" foi removida, uma vez que representava apenas o momento em que o incidente foi registado, não acrescentando informação relevante para o entendimento dos padrões dos incidentes e, por esse motivo, não foi considerada para a análise.

3.4 Tratamento de valores omissos e inválidos

Após a análise individual dos sete valores omissos da variável "Date of the insertion of the SSI", concluiu-se, com base nas variáveis "Date", "Value Date" e "Comment", que se tratava

de incidentes ocorridos em 2022 mas registados apenas em 2023, tendo sido atribuída a data "01-01-2023" a esses valores em falta.

Para a variável "*CLIENT*" alterou-se o seu único valor omissos para "*Unknown*". Tratando-se apenas de um caso, não irá afetar este valor desconhecido de forma relevante para o modelo. De igual forma, para as variáveis "*VALIDATOR*" e "*CREATOR*", os valores omissos foram transformados em "*Unknown*", uma vez que estas variáveis foram utilizadas apenas na Análise Exploratória dos Dados.

A coluna "*Comments*" apresentava 482 valores omissos, correspondendo a 97,37% de dados em falta, pelo que se optou por integrar os conteúdos desta variável na coluna "*Comment*", consolidando a informação numa única variável de comentários no conjunto de dados.

A variável "*Cause*" continha nove valores omissos e quinze registos com o valor "-", totalizando 25 entradas inválidas, que foram analisadas individualmente com base nos conteúdos da variável "*Comment*" e posteriormente corrigidas com a justificativa mais adequada.

Por fim, foram corrigidos alguns valores da variável "*Ccy*", ajustando-os ao formato padronizado, como nos casos em que surgia "USA" em vez de "USD".

3.5 Variáveis adicionadas

Com o intuito de enriquecer a análise exploratória com informação relevante, foram criadas variáveis a partir de variáveis já existentes no *dataset*. Esta abordagem permitiu extrair dimensões complementares dos dados originais, sem introduzir informação externa, potenciando a identificação de padrões subjacentes aos incidentes.

Em particular, a partir da variável "*Date of the insertion of the SSL*", decorreram-se duas novas variáveis, "Dia da Semana" e "Mês". Estas variáveis tiveram como objetivo observar se a distribuição dos incidentes apresentava variações significativas ao longo da semana ou entre diferentes meses do ano.

A introdução destas variáveis temporais possibilitou uma leitura mais detalhada do comportamento dos incidentes ao longo do tempo, permitindo identificar eventuais padrões semanais ou mensais na sua ocorrência e, consequentemente, apoiar uma interpretação mais informada dos resultados obtidos na análise exploratória.

3.6 Tratamento de dados não estruturados para a geração de tópicos

Para dar início ao processo de extração de texto destinado ao modelo LDA foi necessário realizar um conjunto de tratamentos linguísticos preliminares sobre os dados textuais. Este modelo baseia-se numa abordagem probabilística que assume que cada documento é uma combinação de tópicos, e cada tópico é uma distribuição de palavras. No entanto, o LDA não capta relações semânticas entre termos, tratando as palavras como unidades independentes. Por essa razão, é fundamental normalizar e reduzir a variabilidade lexical antes da modelação. Importa notar que este tipo de pré-processamento é específico para modelos baseados em *bag-of-words*, como o LDA, e poderá ser significativamente distinto em abordagens mais recentes como o *BERTopic*, que tiram partido de representações semânticas.

Numa primeira fase, procedeu-se à padronização do conteúdo textual, convertendo todos os caracteres para minúsculas e removendo elementos irrelevantes como URLs e marcas de formatação em HTML. De seguida, aplicou-se um processo de *stemming*, que consiste na redução das palavras às suas formas base ou radicais. Esta etapa tem como objetivo

identificar similaridades linguísticas e reduzir a dimensionalidade do vocabulário, contribuindo para uma análise mais eficiente.

Foi igualmente construído um mapa ou dicionário de sinónimos, destinado a unificar termos redundantes ou semanticamente equivalentes, pois a sua presença dispersa poderia introduzir ruído nos resultados. A definição deste dicionário baseou-se no conhecimento prévio do domínio e numa análise qualitativa extensiva dos comentários associados aos incidentes.

Assim, sempre que surgiram variantes destas expressões, estas passaram a ser tratadas como equivalentes no modelo. Para garantir consistência terminológica, diferentes variantes e abreviações foram normalizadas para uma forma comum. A Tabela 2 apresenta os principais termos identificados e as respetivas formas normalizadas utilizadas no processo de tratamento dos dados.

Em particular, termos como “cut”, “COT” e “cut off” foram normalizados para “CUT-OFF”, uma vez que todos se referem ao mesmo conceito, correspondente à hora de fecho de um determinado mercado. De forma semelhante, “ref” e “team” foram substituídos por “referential”, por se tratar da designação comum atribuída às equipas analisadas. “COSMO” foi substituído por “ALERT”, pois se trata de uma designação interna utilizada para referir a mesma plataforma. Foram também reduzidas redundâncias associadas a abreviações e plurais, como no caso de “value date” substituído por “vd” e “ssis” por “ssi”. A sigla SSI corresponde a *Standard Settlement Instructions* e refere-se às instruções bancárias dos clientes.

Tabela 2: Normalização de termos e expressões no processo de tratamento textual

Variantes Identificadas	Termo Normalizado
ALERT, COSMO	ALERT
BIC, REDACTED	BIC
CUT-OFF, CUT, COT, CUT OFF	CUT-OFF
BENEFICIARY, BENEF	BENEFICIARY
REFERENTIAL, REF, TEAM	REFERENTIAL
VALUE DATE, VD	VD
SSI, SSIS	SSI

Inicialmente, foi utilizada uma lista padrão de *stopwords* da língua inglesa (por exemplo: *a, an, and, are, as, at, be, but, by, for, if, in, into, is, it, no, not, of, on, or, such, that, the, their, then, there, these, they, this, to, was, will, with*, entre outras).

Posteriormente, foram excluídas palavras adicionais, como “*said*”, “*would*”, “*could*”, “*get*”, “*go*” e “*like*”, por não contribuírem de forma relevante para a análise. A palavra “*redacted*” foi igualmente removida, uma vez que apenas surge no contexto da anonimização de dados confidenciais e não contém valor semântico relevante para os fins do modelo.

Durante a fase de testagem do modelo LDA, foram ainda identificadas *stopwords* adicionais que dificultavam a interpretação dos tópicos gerados, por não acrescentarem informação relevante sobre as causas subjacentes aos incidentes analisados. Estas palavras, apesar de não constarem da lista padrão de *stopwords*, revelaram-se recorrentes nos comentários e pouco informativas do ponto de vista semântico. Entre os termos removidos incluem-se referências temporais como por exemplo “*january*”, “*march*”, “*13th*” ou “*20th*”, siglas e nomes próprios como “*fxo*”, “*fxmm*”, “*citi*” ou “*hong kong*”. A exclusão destas palavras teve como

objetivo melhorar a coerência dos tópicos extraídos, reduzindo o ruído textual e permitindo uma representação mais precisa das temáticas centrais associadas aos incidentes. A lista completa de *stopwords* adicionais incluídas encontra-se abaixo:

["january", "23rd", "hong", "kong", "ref", "referential", "team", "fxo", "13th", "11th", "done", "december", "fxmm", "said", "would", "could", "get", "go", "like", "citi", "20th", "xxx", "march", "february", "april", "may", "june", "july", "august", "september", "october", "november"]

3.7 Caracterização dos dados textuais

Após a junção das duas variáveis de comentários, verifica-se que o número médio de palavras por documento é de aproximadamente 37 palavras, sendo que 50% dos documentos apresentam entre 19 e 44 palavras, conforme é possível verificar através da Figura 1.

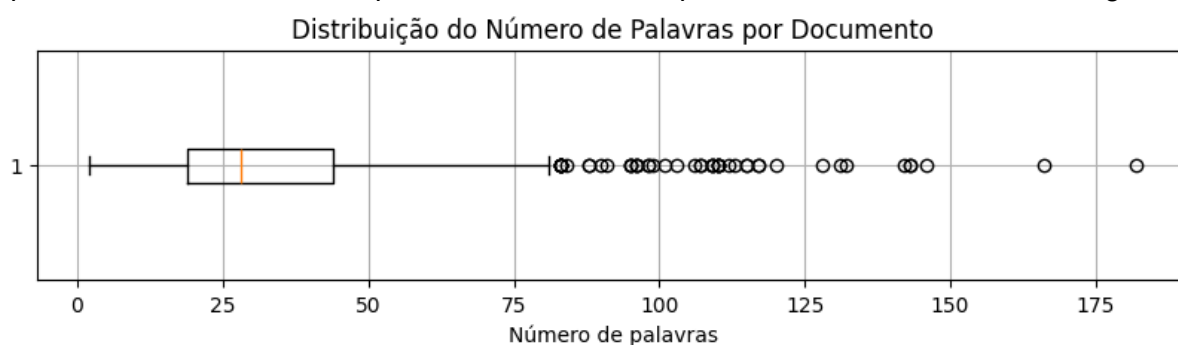


Figura 1: Distribuição do número de palavras por documento

O corpus analisado é composto por 495 documentos, contendo um total de 19.186 palavras e um vocabulário de 1.445 palavras distintas, contando com uma diversidade vocabular (*Type-Token Ratio*) de 0,0753, indicando que 7,5% das palavras utilizadas são distintas, o que sugere que o corpus é relativamente limitado e repetitivo. Esta baixa diversidade é expectável, na medida em que se trata de textos escritos em contexto operacional, onde os documentos tendem a utilizar terminologia padronizada e descrições semelhantes para reportar os incidentes. Entre as palavras distintas, existem 44 que foram utilizadas apenas uma vez.

3.8 Análise exploratória dos dados

A análise exploratória de dados é uma etapa crucial em qualquer projeto de análise de dados, pois permite compreender a estrutura dos dados, identificar padrões, detetar anomalias e testar hipóteses iniciais, que podem guiar análises mais aprofundadas.

Dos 495 incidentes registados, a equipa teve responsabilidade em 299 casos, o que corresponde a 60,4% do total. Em 191 incidentes, não foi atribuída responsabilidade à equipa, e apenas 5 casos já tinham sido reportados previamente com indicação de responsabilidade. A esmagadora maioria dos incidentes não consistiu em pagamentos por engano, representando apenas 12 dos 495. Embora o número seja reduzido, quando comparado com os casos onde não foram feitos pagamentos de forma indevida, os incidentes com pagamentos indevidos tendem a estar associados a perdas significativamente para o banco.

No total, foram identificadas 37 moedas distintas associadas aos incidentes. Entre as dez moedas com maior número de ocorrências encontram-se as seguintes:

Tabela 3: Distribuição das 10 moedas com maior número de incidentes

Moeda	Número de Incidentes
EUR	120

USD	96
JPY	30
AUD	29
GBP	27
HKD	21
CNH	18
CZK	14
CAD	13
THB	13

As dez moedas mais frequentes concentram um total de 381 incidentes, o que representa 76,97% dos registos analisados, evidenciando uma forte concentração em torno de um conjunto reduzido de moedas, como podemos verificar na Tabela 3. Esta concentração justifica uma análise mais aprofundada sobre as suas características. As moedas EUR, USD e GBP são as que têm maior volume de pagamentos registada, e os respetivos *cut-offs* ocorrem tipicamente no final do dia. Este fator sugere que os incidentes associados a estas moedas poderão não estar diretamente relacionados com os prazos de processamento, mas sim com outros elementos do processo, como erros de *setup* ou falhas operacionais.

Por outro lado, moedas como JPY, HKD, CNH e THB são caracterizadas por *cut-offs* antecipados, frequentemente durante a madrugada. Estas restrições temporais exigem que os pagamentos sejam iniciados no início da manhã ou mesmo no dia útil anterior, o que aumenta significativamente o risco de erro, sobretudo quando a equipa ainda não iniciou o seu horário de trabalho. Adicionalmente, moedas como AUD, HKD, CNH e CZK apresentam particularidades técnicas relevantes, relacionadas quer com regras específicas da rede SWIFT, quer com limitações ou comportamentos específicos do *software* utilizado. Embora estas moedas sejam menos comuns no volume total de operações, a sua complexidade pode gerar mais dúvidas ou erros, especialmente entre membros da equipa com menor experiência.

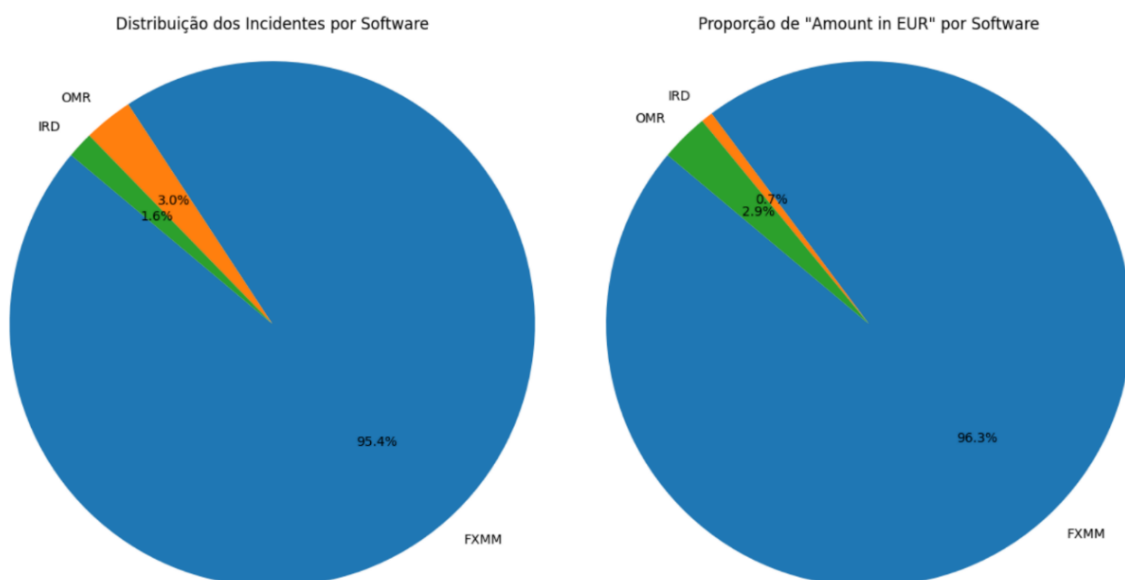


Figura 2: Distribuição da frequência e do valor dos Incidentes por software

Através da Figura 2, percebemos que a maioria dominante dos incidentes registados pertence ao sistema FXMM, com 95,4% de todos os incidentes. Do mesmo modo, este representa a maioria dos montantes associados aos incidentes com 96,3%. Este predomínio justifica-se

pelo volume elevado de operações processadas por este *software*, pela exigência particular relativa a prazos e a casos sensíveis e complexos, como por exemplo TPP, considerados os pagamentos mais sensíveis e com maiores perdas para o banco.

Observa-se ainda que os incidentes registados no sistema IRD, embora correspondam a 1,6% do número total de incidentes, representam apenas 0,7% do valor monetário associado. Esta discrepância sugere que os incidentes ocorridos no IRD tendem a envolver montantes mais reduzidos, quando comparados com os registados noutros sistemas. Por outro lado, apesar de o OMR apenas ter 1,6% dos incidentes, a proporção do valor monetário deste sistema é de quase o dobro da sua frequência, sugerindo que apesar de pouco frequentes, os incidentes relacionados com este sistema envolvem valores monetários maiores.

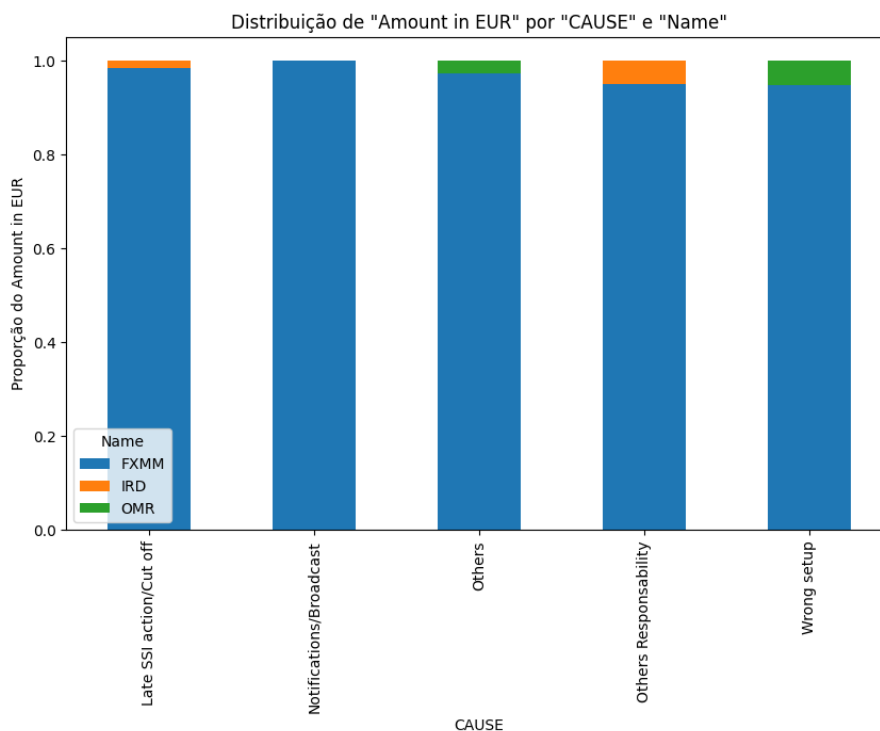


Figura 3: Distribuição proporcional do valor em EUR por causa de incidente, segmentado pelo sistema responsável

A partir do gráfico da Figura 3, percebe-se que a maioria dos incidentes registados no sistema OMR está associada a erros de configuração entre outros, enquanto no caso de IRD, todos os incidentes com responsabilidade da equipa se devem a ações tardias. A inexistência de incidentes causados por atrasos em OMR pode ser explicado pelo reduzido volume de operações processadas, possibilitando uma maior capacidade de resposta atempada. Ainda assim, apesar do baixo volume, os incidentes que foram registados neste sistema demonstram montantes bastante elevados, demonstrando a maior sensibilidade técnica e a maior sensibilidade dos produtos transacionados.

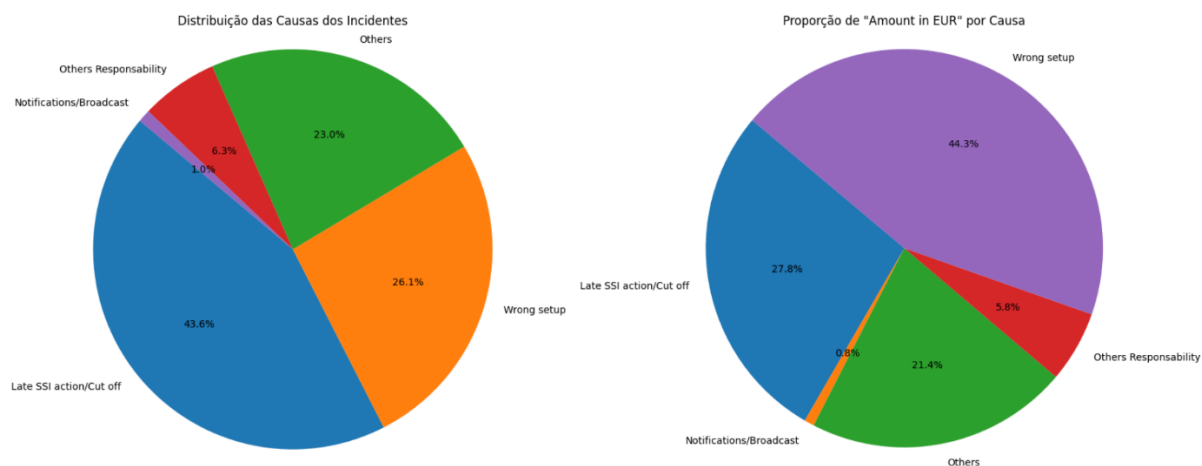


Figura 4: Distribuição e impacto financeiro dos incidentes por causa

Tabela 4: Distribuição de incidentes por causa

Causa	Número de Incidentes
<i>Late SSI action/Cut off</i>	212
<i>Wrong Setup</i>	130
<i>Others</i>	128
<i>Others Responsibility</i>	20
<i>Notifications/Broadcast</i>	5

Através da Figura 4 e da Tabela 4, conseguimos verificar a comparação entre a distribuição da frequência dos incidentes por causa e a respetiva proporção em termos de valor monetário permite identificar diferenças significativas. Embora a categoria "*Late SSI action/Cut off*" represente a causa mais frequente, a causa "*Wrong setup*" assume uma maior relevância quando analisada em função do montante total envolvido. Esta discrepância pode ser explicada pelo facto de os incidentes classificados como "*Wrong setup*" estarem frequentemente associados a perdas substanciais, como pagamentos para contas incorretas, enquanto os incidentes por ação tardia tendem a originar apenas pequenas penalizações, nomeadamente taxas por atraso.

A identificação de tópicos nos incidentes cuja causa é classificada como "*Wrong setup*" reveste-se de particular importância, uma vez que esta categoria está associada aos incidentes com maior impacto financeiro para a organização. Ao aplicar técnicas de modelação de tópicos a este subconjunto específico, torna-se possível identificar padrões linguísticos recorrentes, processos críticos e fragilidades operacionais que originam perdas elevadas. Esta abordagem orientada permite priorizar ações corretivas com base no custo associado, contribuindo para uma mitigação mais eficaz dos riscos operacionais e para uma alocação mais eficiente dos recursos de melhoria contínua.

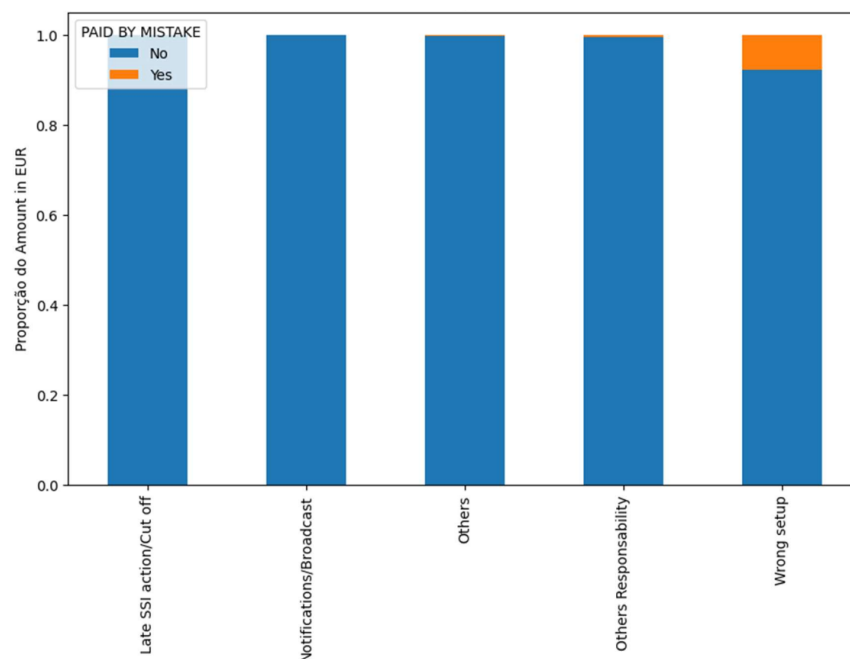


Figura 5: Distribuição proporcional dos montantes pagos ("Amount in EUR") por causa do incidente ("Cause") e ocorrência de pagamento indevido ("Paid by Mistake")

Tabela 5: Distribuição do número de incidentes por causa e ocorrência de pagamento por engano

Cause	PAID BY MISTAKE	
	No	Yes
<i>Late SSI action/Cut</i>	216	0
<i>Notifications/Broadcast</i>	5	0
Others	112	2
Others Responsibility	29	2
Wrong setup	121	8

A análise da distribuição percentual dos pagamentos efetuados por engano, por causa do incidente da Figura 5, complementada pela distribuição do número de incidentes por causa e ocorrência de pagamento por engano apresentada na Tabela 5, mostra que a categoria "*Wrong setup*" é a principal responsável por este tipo de ocorrência. Este resultado reforça a importância de identificar e compreender os fatores subjacentes a esta causa, uma vez que representa a maior proporção de incidentes com impacto financeiro direto.

Por outro lado, não se registaram pagamentos por engano em incidentes associados às causas "*Late SSI action/Cut off*" e "*Notifications/Broadcast*". No primeiro caso, tal deve-se ao facto de o incidente resultar do não processamento atempado do pagamento, ou da sua não execução, impedindo que ocorresse qualquer transferência indevida. No caso dos incidentes classificados como "*Notifications/Broadcast*", a origem está geralmente ligada a processos de atualização de contas. Nestes casos, a principal causa prende-se com a falta de atualização atempada das instruções por parte dos clientes ou das equipas internas. Contudo, estas instruções referem-se, na maioria das vezes, a contas inativas, pelo que não são gerados pagamentos, evitando-se, assim, a ocorrência de pagamentos incorretos.

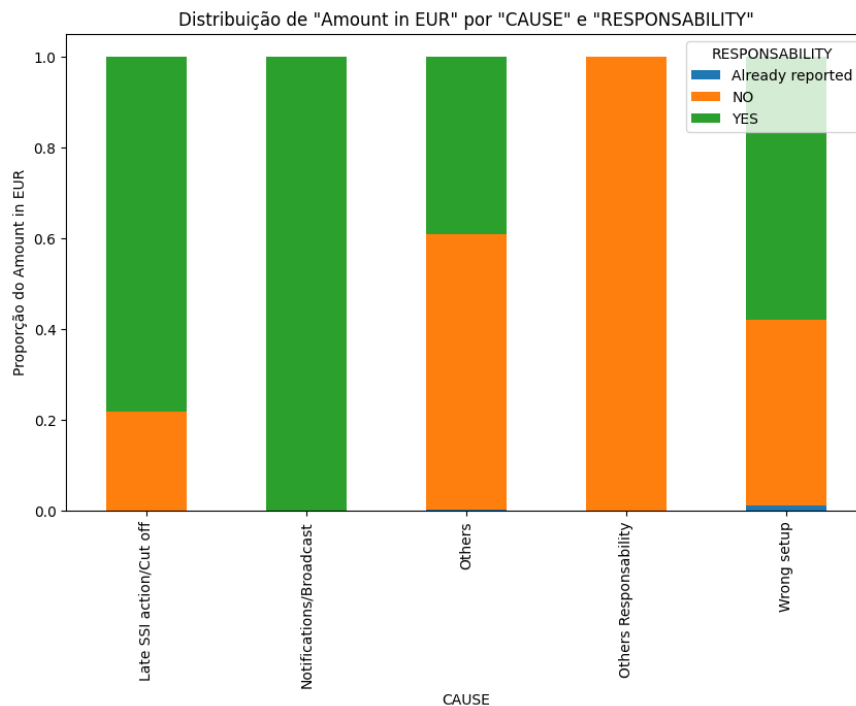


Figura 6: Distribuição proporcional do montante em EUR por causa do incidente e por responsabilidade

Dos incidentes registrados em 2023, 299 tiveram responsabilidade atribuída às equipes deste departamento, o que representa 60,4% do total. Entre estes, a principal causa identificada foi a ação tardia em relação ao *cut-off*, resultando em atrasos nos pagamentos, correspondendo a 56,19% dos incidentes registrados. Seguem-se os erros de configuração (*wrong setup*), com 89 incidentes, e a categoria "*others*", com 37 incidentes, que em conjunto representam 42,14% dos casos com responsabilidade atribuída à equipe. Estas duas últimas causas, pela sua diversidade e natureza menos específica, são particularmente relevantes para a aplicação de técnicas de identificação de tópicos, permitindo aprofundar a compreensão das causas subjacentes aos incidentes associados. Conforme a Figura 6, observa-se que a categoria "Notifications/Broadcast" apresenta a totalidade do montante associada a incidentes em que a responsabilidade foi assumida. A categoria "Wrong setup" evidencia uma distribuição entre ambas as classificações, embora com predominância de incidentes em que a responsabilidade foi assumida.

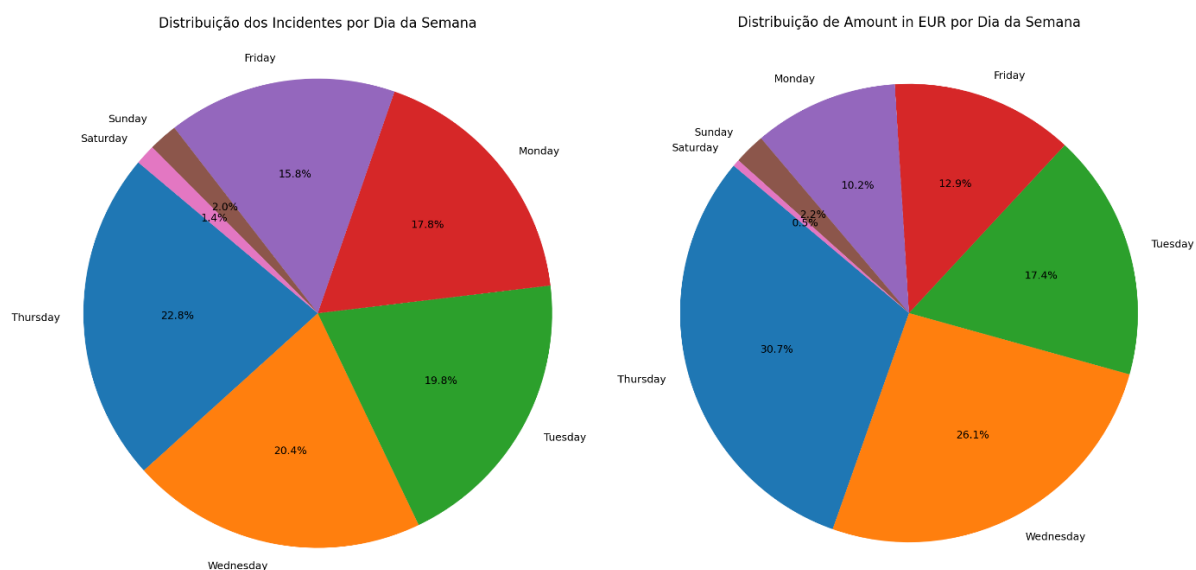


Figura 7: Distribuição de incidentes por dias da semana e valor dos incidentes em euros por dia da semana

A análise dos gráficos da Figura 7 permite observar que a distribuição dos incidentes se mantém relativamente estável ao longo dos dias úteis da semana. Como seria de esperar, o número de incidentes ao fim de semana é bastante reduzido, uma vez que não há atividade operacional durante esse período e são raros os pagamentos agendados também para esse período. No que respeita à distribuição do valor dos incidentes em euros por dia da semana, verifica-se uma maior concentração de montantes mais elevados entre quarta e quinta-feira.

Tabela 6: Distribuição dos incidentes por mês em frequência e valor monetário

Mês	Incidentes (%)	Amount EUR (%)
January	7,70%	6,00%
February	12,10%	12,70%
March	7,30%	7,40%
April	9,50%	6,90%
May	5,90%	2,90%
June	11,50%	9,30%
July	5,90%	9,20%
August	3,60%	1,40%
September	11,30%	8,40%
October	10,50%	17,80%
November	5,70%	8,90%
December	9,10%	9,40%

A análise da Tabela 6 revela que os meses de fevereiro, abril, junho, setembro, outubro e dezembro concentram a maior parte dos incidentes, com uma distribuição individual entre 9% e 12% do total. Em conjunto, estes meses representam aproximadamente 64% dos incidentes registados no *dataset*. Por outro lado, agosto apresenta o número mais reduzido de incidentes, com apenas 3,6% dos registos, o que se justifica pela significativa diminuição da atividade operacional durante este período, tradicionalmente associado a férias e a um menor volume de pagamentos.

A análise da distribuição do valor dos incidentes em euros por mês revela que outubro, fevereiro e dezembro concentram os montantes mais elevados, representando respetivamente 17,8%, 12,7% e 9,4% do total. Esta distribuição financeira está em linha com a frequência de incidentes anteriormente identificada, indicando que os meses com maior número de ocorrências tendem também a concentrar valores monetários mais significativos.

Destaca-se ainda o mês de agosto, que além de apresentar o número mais reduzido de incidentes (3,6% dos registos), representa apenas 1,4% do valor total em euros. Este comportamento confirma a redução acentuada da atividade operacional durante o período de férias, tanto em volume como em impacto financeiro.

3.9 Considerações finais

A análise exploratória permitiu compreender em profundidade a estrutura dos dados e os padrões operacionais associados aos incidentes registados em 2023. Embora a causa mais recorrente tenha sido “*Late SSI action/Cut off*”, foi a categoria “*Wrong Setup*” que se destacou pelo seu impacto financeiro, concentrando a maioria dos pagamentos efetuados por engano. Esta constatação evidencia a necessidade de uma análise mais apertada desta causa, que se revela crítica para a mitigação do risco operacional.

Adicionalmente, observaram-se fatores contextuais relevantes, como a maior incidência de incidentes em moedas com *cut-offs* antecipados (e.g., JPY, CNH, THB) e a concentração de ocorrências no sistema FXMM, o que aponta para constrangimentos técnicos e processuais específicos. Com base nestes *insights*, a próxima etapa do estudo consistirá na construção de um modelo de *text mining*, assente na variável “*Comment*”, com o objetivo de identificar tópicos recorrentes nos incidentes de “*Wrong Setup*”. A aplicação de técnicas de modelação de tópicos permitirá extrair padrões linguísticos e operacionais, contribuindo para uma identificação mais precisa das causas subjacentes e para uma atuação corretiva mais eficaz.

Esta página foi intencionalmente deixada em branco

4 Modelação

Para a identificação de tópicos latentes, foram considerados 126 incidentes em que a responsabilidade foi atribuída às equipas que reportaram o incidente e cuja causa se enquadrava nas categorias "Wrong Setup" ou "Others". A causa "Late SSI action/Cut off" foi excluída da análise por se tratar de uma razão autoexplicativa, cuja inclusão poderia introduzir ruído nos modelos de identificação de tópicos. A variável "Others Responsibility" também não foi considerada, uma vez que não existe responsabilidade atribuída à equipa que reportou o incidente, o que limita significativamente o seu conhecimento sobre a origem do problema. Por fim, os cinco incidentes classificados como "Notifications/Broadcast" foram analisados individualmente, tendo-se verificado que, em todos os casos, a causa subjacente foi uma ação tardia por parte dos colaboradores. Assim, estes foram reclassificados como pertencentes à categoria "Late SSI action/Cut off" e, por conseguinte, também excluídos da análise pelos mesmos motivos acima referidos.

Nesta fase procede-se à aplicação de várias técnicas de *text mining* para a extração de tópicos latentes relevantes para a identificação das causas dos incidentes financeiros. Para o presente projeto foram utilizados três modelos distintos identificados durante a revisão da literatura, cada um deles com diferentes abordagens técnicas: um modelo probabilístico, o LDA, que identifica conjuntos de palavras frequentemente associadas e as organiza em tópicos; um modelo baseado em representações vetoriais, o BERTopic, capaz de detetar nuances semânticas na geração de tópicos e, por fim, foi implementada uma abordagem que combina o BERTopic com o uso de *Large Language Models* (LLM), com o objetivo de auxiliar na interpretação e nomeação dos tópicos gerados.

Nesta configuração, o LLM é utilizado para analisar o contexto semântico das palavras mais relevantes de cada tópico, permitindo uma rotulagem mais precisa e contextualizada, complementando assim a representação vetorial produzida pelo BERTopic.

4.1 Latent Dirichlet Allocation

O *Latent Dirichlet Allocation* é um modelo probabilístico de geração de tópicos que permite identificar estruturas latentes em coleções de documentos. O princípio fundamental assenta na suposição de que cada documento é representado por uma mistura de tópicos latentes, e que cada tópico corresponde a uma distribuição probabilística sobre o vocabulário (Blei, Ng e Jordan 2003). Desta forma, a técnica permite identificar padrões de coocorrência de palavras e revelar estruturas temáticas subjacentes em grandes volumes de texto não estruturado.

O modelo é formalizado como um modelo hierárquico bayesiano de três níveis, no qual o *corpus*, os documentos, os tópicos e as palavras se encontram interligados. No primeiro nível, cada tópico é descrito por uma distribuição de palavras extraída de uma distribuição de *Dirichlet*. No segundo nível, cada documento possui a sua própria mistura de tópicos, também gerada a partir de uma distribuição de *Dirichlet*, a qual define a proporção com que esses tópicos aparecem no texto. Finalmente, no terceiro nível, cada palavra de um documento é referenciada escolhendo-se um tópico segundo a distribuição do documento e, em seguida, escolhendo aleatoriamente uma palavra da distribuição desse tópico. Esta estrutura probabilística hierárquica permite capturar simultaneamente as regularidades globais do *corpus* e as especificidades de cada documento (Blei, Ng & Jordan, 2003).

Do ponto de vista matemático, o LDA pode assim ser descrito: para cada documento d , a distribuição de tópicos θ_d é extraída de uma *dirichlet* com *hiper-parâmetro* α , sendo θ_d o vetor de proporções de tópicos associados ao documento d e o parâmetro α que controla a

dispersão dessa distribuição. Para cada palavra W_{dn} , escolhe-se um tópico $Z_{dn} \sim \text{Multinomial}(\theta_d)$, sendo a palavra então gerada de acordo com a distribuição $p(W_{dn}/Z_{dn}, \beta)$, onde β representa as distribuições de palavras associadas a cada tópico, Z_{dn} correspondendo à variável latente que indica o tópico atribuído à n -ésima palavra do documento d e W_{dn} à palavra observada nessa posição. A probabilidade global de um documento é, assim, obtida através da integração sobre todas as possíveis distribuições de tópicos latentes, o que permite modelar explicitamente a incerteza associada à atribuição de tópicos e confere ao LDA uma base estatística sólida para a descoberta de temas latentes.

A avaliação da qualidade dos tópicos gerados pode ser realizada através de diferentes métricas, de entre as quais se destaca a *coerência*. Esta métrica quantifica o grau de similaridade semântica entre as palavras mais representativas de cada tópico, permitindo aferir o nível de consistência interna dos temas identificados. Modelos com valores elevados de coerência tendem a produzir tópicos cujas palavras apresentam maior proximidade em termos de significado e contexto, o que se traduz numa maior capacidade de interpretação e numa representação mais fiável das relações latentes no *corpus* analisado.

Apesar da utilidade da coerência enquanto métrica extrínseca, esta desempenha apenas um papel complementar neste projeto, devido à sua natureza não supervisionada, o que faz com que a validação dos tópicos dependa sobretudo de uma avaliação qualitativa, baseada no conhecimento de negócio para poder interpretar os resultados e nomeá-los de forma adequada, garantindo relevância prática para o contexto operacional analisado.

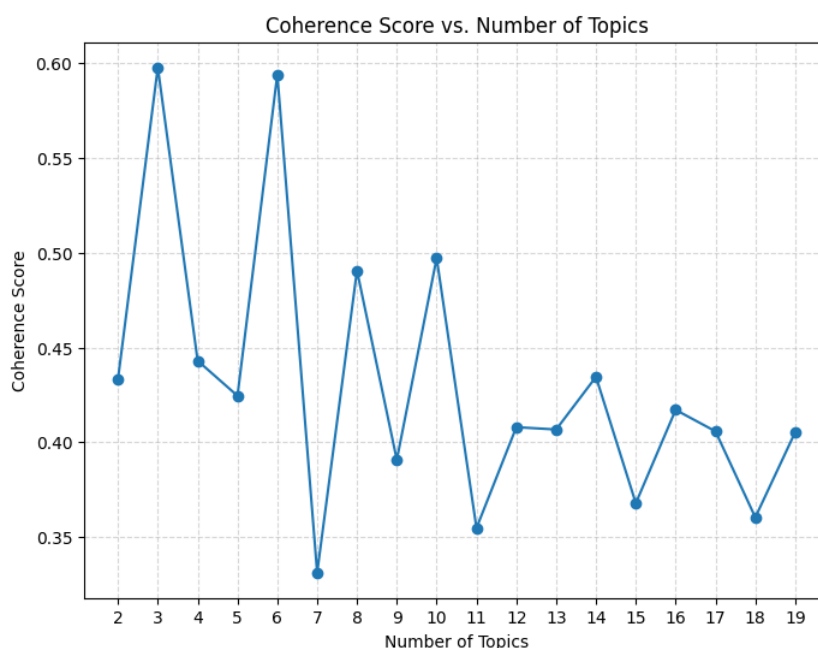


Figura 8: Variação da coerência em função do número de tópicos no modelo LDA

A Figura 8 apresenta a variação da coerência em função do número de tópicos gerados pelo modelo, permitindo seleccionar diferentes configurações para posterior análise. Verifica-se que o valor máximo de coerência é alcançado quando o modelo é parametrizado com três tópicos, atingindo 0,60 numa escala de 0 a 1, sendo a configuração com seis tópicos a que mais se aproxima deste desempenho, com valores apenas marginalmente inferiores. Os valores seguintes mais elevados correspondem às configurações com 8 e 10 tópicos, situando-se entre 0,45 e 0,50, isto é, relativamente inferiores aos modelos com três e seis tópicos, mas ainda assim relevantes.

Para este projeto, serão realizadas experiências com modelos de 3, 6, 8 e 10 tópicos, partindo do pressuposto de que, embora as soluções com maior número de tópicos apresentem coerência inferior, estas podem contribuir para a identificação de tópicos menos frequentes, mas igualmente significativos, que tenderiam a ser agregados num modelo mais restritivo. Tal abordagem, contudo, implica um compromisso, dado que a proliferação de tópicos aumenta a probabilidade de sobreposição temática e reduz a clareza da separação entre eles.

A representação visual dos tópicos será realizada através de diferentes abordagens complementares: (1) uma nuvem de palavras, em que o tamanho de cada termo é proporcional à sua relevância no tópico; (2) uma rede de conceitos, que evidencia as ligações entre os principais termos e (3) a listagem das palavras mais representativas de cada tópico acompanhadas do respetivo grau de importância. A designação de cada tópico será proposta no presente estudo, com base no contexto e no conhecimento de negócio.

Após a extração e interpretação dos tópicos, proceder-se-á igualmente à análise do valor de coerência, bem como à avaliação da distribuição dos tópicos pelos documentos e à da sua separação através da análise de um mapa de componentes principais.

4.1.1 Abordagem com 3 tópicos

A configuração do modelo LDA com três tópicos foi selecionada como ponto de partida da análise, uma vez que apresentou o valor mais elevado de coerência, conforme evidenciado na Figura 8. Esta escolha justifica-se pela capacidade acrescida desta parametrização em produzir tópicos mais consistentes e interpretáveis, assegurando uma melhor separação semântica entre os grupos de termos identificados. A análise com três tópicos permite, assim, explorar de forma clara as principais dimensões latentes presentes nos incidentes operacionais, oferecendo uma visão inicial sobre os padrões mais relevantes no *corpus*.



Figura 9: Nuvens de palavras representativas dos três tópicos gerados pelo modelo LDA

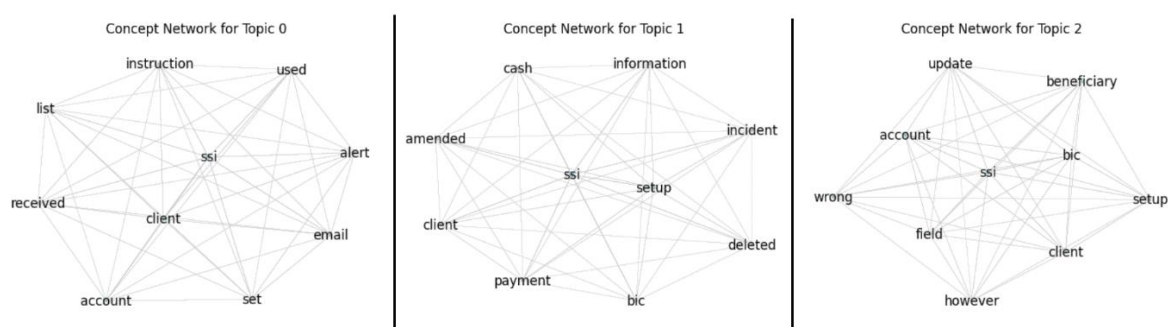


Figura 10: Redes de conceitos associadas aos três tópicos gerados pelo modelo.

A partir das nuvens de palavras da Figura 9, conseguimos verificar quais as palavras predominantes em cada tópico, estando estas organizadas de uma forma em que as palavras mais relevantes aparecem maiores na nuvem de palavras, enquanto as palavras menos comuns aparecem menores. Por sua vez, a Figura 10 permite aferir a relação dos principais

conceitos entre si, sendo que os conceitos que surgem no meio do espectro são aqueles que têm mais relações com os restantes.

Para além disso, extraiu-se uma lista com as 10 palavras predominantes de cada tópico e as suas respetivas importâncias dentro do tópico em valor percentual:

- **Topic 1:** $0.054 \cdot \text{'ssi'}$ + $0.037 \cdot \text{'client'}$ + $0.024 \cdot \text{'email'}$ + $0.022 \cdot \text{'alert'}$ + $0.022 \cdot \text{'received'}$ + $0.020 \cdot \text{'set'}$ + $0.020 \cdot \text{'instruction'}$ + $0.017 \cdot \text{'used'}$ + $0.016 \cdot \text{'account'}$ + $0.016 \cdot \text{'list'}$
- **Topic 2:** $0.067 \cdot \text{'ssi'}$ + $0.033 \cdot \text{'setup'}$ + $0.017 \cdot \text{'payment'}$ + $0.016 \cdot \text{'client'}$ + $0.014 \cdot \text{'amended'}$ + $0.014 \cdot \text{'bic'}$ + $0.012 \cdot \text{'incident'}$ + $0.011 \cdot \text{'cash'}$ + $0.011 \cdot \text{'deleted'}$ + $0.011 \cdot \text{'information'}$
- **Topic 3:** $0.045 \cdot \text{'ssi'}$ + $0.035 \cdot \text{'bic'}$ + $0.030 \cdot \text{'account'}$ + $0.025 \cdot \text{'client'}$ + $0.024 \cdot \text{'field'}$ + $0.018 \cdot \text{'setup'}$ + $0.016 \cdot \text{'beneficiary'}$ + $0.015 \cdot \text{'wrong'}$ + $0.013 \cdot \text{'update'}$ + $0.013 \cdot \text{'however'}$

A interpretação qualitativa permitiu a nomeação dos tópicos, recorrendo ao conhecimento de ação para a sua contextualização.

- Tópico 1 – Incidentes Relacionados com a plataforma ALERT
- Tópico 2 – *Setups* incorretos por conta de um BIC apagado do sistema
- Tópico 3 – Beneficiário errado

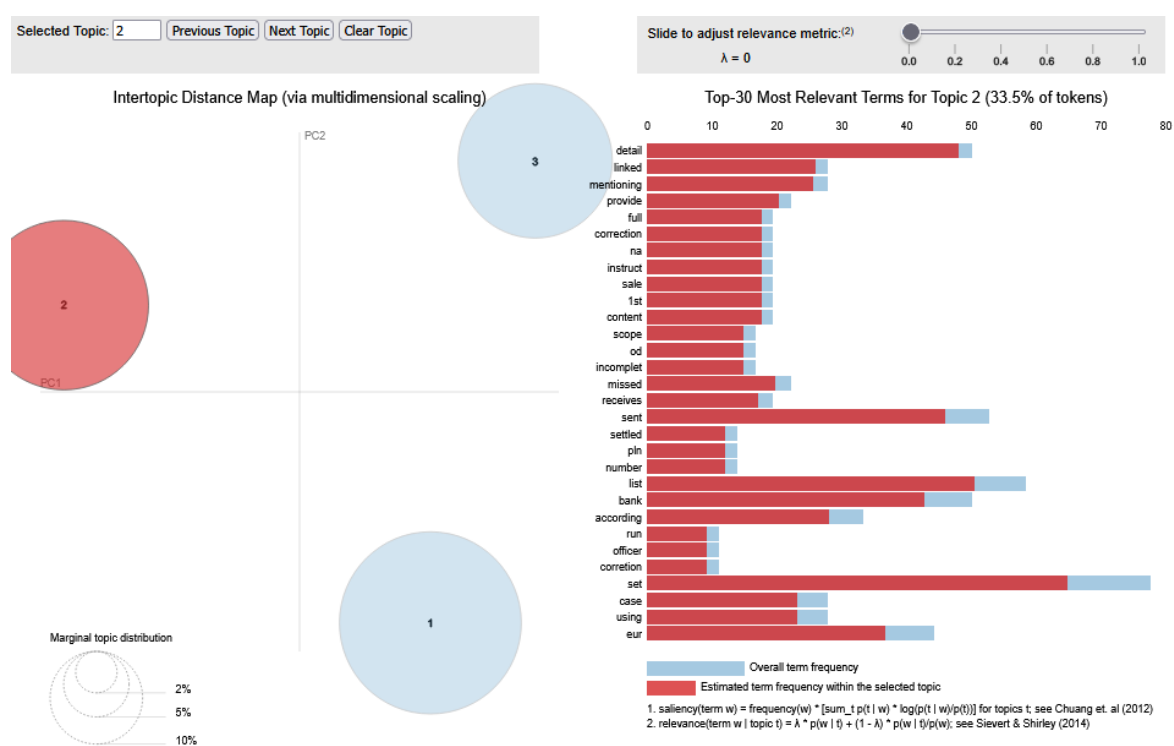


Figura 11: Mapa de Componentes Principais e termos mais relevantes

A biblioteca *pyLDAvis* permite gerar um *dashboard* interativo que apresenta a dimensão dos tópicos, isto é, as 30 palavras mais relevantes associadas a cada um e a sua separação num espaço de componentes principais. A análise da Figura 11 mostra que os tópicos se encontram bem delimitados entre si, não se verificando sobreposições temáticas. Observa-se ainda que os tópicos possuem dimensões semelhantes, sem disparidades relevantes no seu

tamanho relativo. Esta constatação é corroborada pela Figura 12, que ilustra a distribuição dos documentos pelos tópicos dominantes.

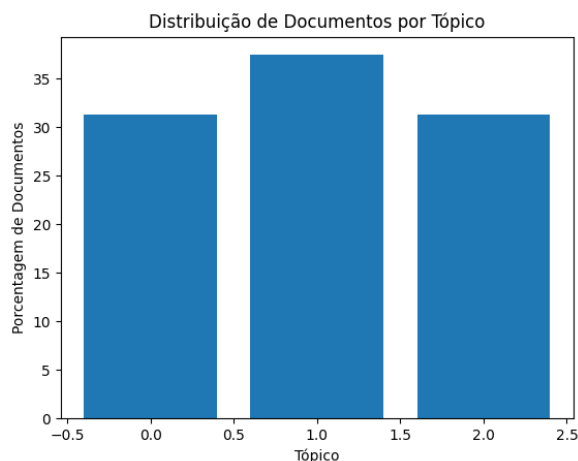


Figura 12: Distribuição percentual de documentos por tópico no modelo LDA com três tópicos

No que respeita à avaliação extrínseca do modelo, a métrica de coerência registou um valor de 0,60. Embora valores mais elevados de coerência indiquem, em geral, tópicos mais consistentes e interpretáveis, importa sublinhar que este indicador deve ser analisado de forma relativa e não absoluta. Assim, o resultado obtido sugere um desempenho aceitável do modelo, refletindo-se em tópicos consistentes e fáceis de interpretar.

4.1.2 Abordagem com 6 tópicos

Para a presente abordagem, bem como para as subseqüentes, serão apresentados apenas os valores de importância de cada tópico, os nomes atribuídos e o respetivo mapa de dimensões, que permite representar visualmente a separação temática. A informação complementar encontra-se disponibilizada em anexo, para efeitos de consulta.

- **Topic 1:** $0.059 \cdot \text{'ssi'} + 0.035 \cdot \text{'setup'} + 0.031 \cdot \text{'client'} + 0.017 \cdot \text{'alert'} + 0.016 \cdot \text{'email'} + 0.015 \cdot \text{'data'} + 0.013 \cdot \text{'set'} + 0.013 \cdot \text{'however'} + 0.013 \cdot \text{'default'} + 0.011 \cdot \text{'place'}$
- **Topic 2:** $0.034 \cdot \text{'update'} + 0.023 \cdot \text{'beneficiary'} + 0.023 \cdot \text{'ssi'} + 0.018 \cdot \text{'however'} + 0.018 \cdot \text{'wrong'} + 0.018 \cdot \text{'changed'} + 0.018 \cdot \text{'currency'} + 0.018 \cdot \text{'incident'} + 0.012 \cdot \text{'name'} + 0.012 \cdot \text{'undisclosed'}$
- **Topic 3:** $0.055 \cdot \text{'ssi'} + 0.047 \cdot \text{'account'} + 0.026 \cdot \text{'field'} + 0.023 \cdot \text{'eur'} + 0.022 \cdot \text{'set'} + 0.021 \cdot \text{'alert'} + 0.020 \cdot \text{'received'} + 0.020 \cdot \text{'detail'} + 0.018 \cdot \text{'wrongly'} + 0.016 \cdot \text{'email'}$
- **Topic 4:** $0.081 \cdot \text{'ssi'} + 0.041 \cdot \text{'client'} + 0.023 \cdot \text{'received'} + 0.021 \cdot \text{'payment'} + 0.020 \cdot \text{'list'} + 0.019 \cdot \text{'email'} + 0.018 \cdot \text{'deleted'} + 0.018 \cdot \text{'bank'} + 0.017 \cdot \text{'updated'} + 0.017 \cdot \text{'setup'}$
- **Topic 5:** $0.068 \cdot \text{'bic'} + 0.039 \cdot \text{'client'} + 0.035 \cdot \text{'ssi'} + 0.020 \cdot \text{'field'} + 0.017 \cdot \text{'wrong'} + 0.016 \cdot \text{'account'} + 0.016 \cdot \text{'setup'} + 0.015 \cdot \text{'loan'} + 0.014 \cdot \text{'alert'} + 0.014 \cdot \text{'missing'}$
- **Topic 6:** $0.049 \cdot \text{'ssi'} + 0.028 \cdot \text{'setup'} + 0.028 \cdot \text{'bic'} + 0.020 \cdot \text{'information'} + 0.019 \cdot \text{'cash'} + 0.017 \cdot \text{'currency'} + 0.016 \cdot \text{'correct'} + 0.015 \cdot \text{'change'} + 0.015 \cdot \text{'fx'} + 0.014 \cdot \text{'instruction'}$

O aumento do número de tópicos permitiu identificar novas potenciais causas de incidentes, embora à custa de alguma sobreposição temática entre eles. A maioria dos tópicos pode ser interpretada de forma relativamente direta, a partir da análise das palavras-chave mais relevantes e das suas relações. A principal exceção corresponde ao Tópico 4, cujos termos apresentam maior ambiguidade semântica. Palavras como “Deleted”, “Bank” e “Updated” podem remeter para diferentes cenários, nomeadamente instruções eliminadas, registos bancários apagados ou atualizados no sistema. Tendo em conta esta análise, os nomes atribuídos aos tópicos foram os seguintes:

- Tópico 1 – Notificações em *Alert*
- Tópico 2 – Beneficiário errado/*Undisclosed Client*
- Tópico 3 – Detalhes de alert mal atualizados
- Tópico 4 – Instruções Apagadas
- Tópico 5 – *BIC* errado, relacionado com pagamentos de *Loans*
- Tópico 6 – Incidente relacionado com a migração

No modelo com seis tópicos, a métrica de coerência registou um valor de 0,59. Este resultado indica um desempenho ligeiramente inferior ao observado em configurações com menor número de tópicos, sugerindo que o aumento da granularidade temática pode ter introduzido maior dispersão semântica e reduzido a consistência dos tópicos, o que reforça as conclusões obtidas na análise qualitativa prévia, adicionando um número de causas identificadas por conta de uma redução marginal na consistência dos tópicos.

Esta solução também mostra uma das principais limitações de modelos probabilísticos como o LDA, relacionadas com a ambiguidade semântica referida sobre o tópico 4, onde a análise dos resultados permite obter diferentes interpretações o conjunto de incidentes representados no tópico.

4.1.3 Abordagem com 8 tópicos

A presente solução com oito tópicos revelou uma sobreposição temática significativamente superior à observada nos modelos anteriores, com vários tópicos a apresentarem interseções entre si, conforme é possível verificar no Anexo B.

- **Tópico 1 – Nome do beneficiário errado, cliente undisclosed:** $0.083^{*}ssi' + 0.035^{*}amended' + 0.032^{*}payment' + 0.026^{*}deleted' + 0.020^{*}setup' + 0.020^{*}benf' + 0.019^{*}incident' + 0.016^{*}change' + 0.016^{*}undisclosed' + 0.014^{*}created'$
- **Tópico 2 – Alert não está em linha com o Bank Check:** $0.046^{*}used' + 0.045^{*}ssi' + 0.033^{*}client' + 0.032^{*}alert' + 0.022^{*}one' + 0.020^{*}account' + 0.019^{*}check' + 0.015^{*}bank' + 0.014^{*}instruction' + 0.014^{*}created'$
- **Tópico 3 – IBAN incorreto:** $0.053^{*}account' + 0.048^{*}ssi' + 0.040^{*}missing' + 0.027^{*}client' + 0.027^{*}incorrect' + 0.027^{*}iban' + 0.020^{*}digit' + 0.014^{*}payment' + 0.014^{*}used' + 0.014^{*}intermediary'$
- **Tópico 4 – Incidente de OMR: Atlas account errada:** $0.024^{*}set' + 0.018^{*}ssi' + 0.017^{*}however' + 0.016^{*}account' + 0.013^{*}updated' + 0.013^{*}info' + 0.013^{*}according' + 0.013^{*}procedure' + 0.013^{*}atlas' + 0.013^{*}note'$
- **Tópico 5 – Campo errado:** $0.083^{*}ssi' + 0.058^{*}client' + 0.047^{*}account' + 0.046^{*}email' + 0.027^{*}received' + 0.025^{*}field' + 0.023^{*}list' + 0.020^{*}sent' + 0.019^{*}wrongly' + 0.018^{*}instruction'$
- **Tópico 6 – Incidente de OMR: Loans:** $0.081^{*}bic' + 0.038^{*}data' + 0.032^{*}setup' + 0.022^{*}client' + 0.020^{*}field' + 0.019^{*}loan' + 0.019^{*}wrong' + 0.014^{*}receives' + 0.013^{*}static' + 0.013^{*}email'$
- **Tópico 7 – Incidente relacionado com um update:** $0.061^{*}ssi' + 0.037^{*}setup' + 0.019^{*}bic' + 0.017^{*}deal' + 0.017^{*}alert' + 0.016^{*}information' + 0.016^{*}instruction' + 0.015^{*}received' + 0.014^{*}however' + 0.014^{*}date'$
- **Tópico 8 – Incidente relacionado com Alert details mal linkados:** $0.054^{*}ssi' + 0.033^{*}client' + 0.030^{*}detail' + 0.029^{*}alert' + 0.022^{*}setup' + 0.022^{*}payment' + 0.022^{*}linked' + 0.017^{*}sett' + 0.016^{*}incident' + 0.014^{*}intermediary'$

Esta sobreposição dificultou a interpretação de alguns tópicos, considerando que se trata de um modelo puramente estatístico, que não considera a estrutura semântica dos documentos. Em particular, destacam-se os Tópicos 2 e 7, que apresentam semelhanças com o Tópico 4 da abordagem com seis tópicos, nos quais a ordem e combinação dos termos podem conduzir a múltiplas interpretações plausíveis sobre a causa do incidente. No caso do segundo tópico, a predominância dos termos '*Alert*', '*Bank*' e '*Check*' sugere a possibilidade de uma atualização na plataforma *Alert* que não estaria alinhada com a informação disponível na *Bank Check*, plataforma utilizada para a verificação de intermediários bancários em determinadas moedas. Já o Tópico 7 mostra ser mais difícil de interpretar, dada a ausência de termos

suficientemente específicos que permitam compreender a origem da falha. Palavras como *'Alert'*, *'Deal'*, *'Received'* e *'However'* apresentam carácter demasiado genérico para permitir um entendimento claro dos incidentes associados a este tópico.

Apesar destas limitações, esta abordagem permitiu identificar incidentes que não haviam sido detetados pelos modelos anteriores. O Tópico 1 aborda situações relacionadas com clientes, geralmente fundos, que preferem não ter o seu nome divulgado no campo do beneficiário, solicitando a utilização de uma designação alternativa. Os Tópicos 4 e 6 evidenciam, por sua vez, incidentes registados num sistema de maior sensibilidade, associados a operações de empréstimo. Ainda que tenha permitido identificar novos padrões, a maioria dos tópicos adicionais revelou-se demasiado semelhante ou, em alguns casos, redundante face aos já identificados nas abordagens precedentes.

No modelo com oito tópicos, a métrica de coerência apresentou um valor de 0,48, evidenciando um desempenho inferior face aos modelos anteriores. Este resultado sugere menor consistência semântica entre as palavras que compõem cada tópico, refletindo a maior sobreposição temática observada e a dificuldade acrescida na interpretação de alguns grupos de termos. Ainda assim, esta abordagem permitiu identificar novas causas de incidentes, relacionadas nomeadamente com dados sensíveis, instruções não divulgadas e operações de empréstimo, embora com um custo claro em termos de coerência e clareza temática.

4.1.4 Abordagem com 10 tópicos

A última abordagem, com dez tópicos, revelou uma elevada sobreposição temática, com vários tópicos a ocuparem praticamente o mesmo espaço no mapa de dimensões, conforme ilustrado no Anexo C. Esta sobreposição resultou não só em tópicos repetidos, mas também em tópicos de difícil interpretação.

- **Tópico 1 – Utilização incorreta de contas ou intermediários:** $0.051^{*}ssi' + 0.048^{*}used' + 0.041^{*}account' + 0.031^{*}client' + 0.026^{*}alert' + 0.026^{*}intermediary' + 0.025^{*}bank' + 0.022^{*}one' + 0.019^{*}iban' + 0.016^{*}created'$
- **Tópico 2 – Setup desatualizado, notificação não processada:** $0.043^{*}setup' + 0.034^{*}ssi' + 0.024^{*}place' + 0.022^{*}since' + 0.020^{*}notification' + 0.019^{*}updated' + 0.016^{*}client' + 0.015^{*}incident' + 0.015^{*}missing' + 0.011^{*}requested'$
- **Tópico 3 – Setup feito com instruções erradas, corrigido após o pagamento:** $0.049^{*}ssi' + 0.031^{*}eur' + 0.026^{*}client' + 0.024^{*}set' + 0.018^{*}deal' + 0.018^{*}received' + 0.017^{*}instruction' + 0.015^{*}setup' + 0.015^{*}amended' + 0.015^{*}bnp'$
- **Tópico 4 – Instruções antigas em sistema, não atualizadas por Alert:** $0.030^{*}data' + 0.030^{*}setup' + 0.024^{*}update' + 0.022^{*}according' + 0.022^{*}instruction' + 0.020^{*}set' + 0.020^{*}old' + 0.016^{*}ssi' + 0.015^{*}email' + 0.014^{*}needed'$
- **Tópico 5 – Erro relacionado com autodébito:** $0.077^{*}ssi' + 0.033^{*}email' + 0.029^{*}client' + 0.025^{*}incident' + 0.021^{*}settlement' + 0.017^{*}correct' + 0.017^{*}sent' + 0.016^{*}wrongly' + 0.016^{*}mismatch' + 0.012^{*}autodebit'$
- **Tópico 6 – Instruções apagadas, erro relacionado com o nome do beneficiário:** $0.098^{*}ssi' + 0.037^{*}client' + 0.026^{*}payment' + 0.023^{*}received' + 0.021^{*}deleted' + 0.020^{*}list' + 0.019^{*}email' + 0.016^{*}updated' + 0.016^{*}correctly' + 0.016^{*}benf'$
- **Tópico 7 – Incidente relacionado com Alert e campo errado:** $0.072^{*}ssi' + 0.068^{*}client' + 0.022^{*}setup' + 0.022^{*}information' + 0.016^{*}created' + 0.016^{*}email' + 0.016^{*}field' + 0.016^{*}alert' + 0.013^{*}correct' + 0.013^{*}default'$
- **Tópico 8 – Erros em setups de Loans:** $0.122^{*}bic' + 0.040^{*}loan' + 0.037^{*}setup' + 0.032^{*}ssi' + 0.031^{*}wrong' + 0.022^{*}default' + 0.016^{*}client' + 0.012^{*}updated' + 0.012^{*}ss' + 0.012^{*}asked'$
- **Tópico 9 – Alert details mal linkados:** $0.057^{*}ssi' + 0.043^{*}alert' + 0.031^{*}bic' + 0.030^{*}setup' + 0.025^{*}detail' + 0.022^{*}instruction' + 0.022^{*}cash' + 0.020^{*}information' + 0.018^{*}however' + 0.016^{*}linked'$
- **Tópico 10 – Erro na conta do beneficiário final:** $0.159^{*}account' + 0.093^{*}field' + 0.066^{*}beneficiary' + 0.050^{*}final' + 0.046^{*}wrongly' + 0.041^{*}instead' + 0.037^{*}inserted' + 0.037^{*}missing' + 0.025^{*}incorrect' + 0.019^{*}digit'$

Relativamente aos tópicos repetidos, estes podem ser agrupados em dois conjuntos temáticos principais. O primeiro está associado a processos de atualização e inclui os Tópicos 2, 3, 4 e 6, que partilham termos como *'setup'*, *'updated'*, *'instruction'*, *'received'*, *'deleted'*, *'email'* e *'client'*. O segundo grupo está relacionado com a plataforma *Alert*, abrangendo os Tópicos 1, 7 e 9, que apresentam em comum termos como *'Alert'*, *'Details'*, *'Information'* e *'Field'*.

De entre os tópicos de interpretação mais ambígua destaca-se o Tópico 3, cujos termos são excessivamente genéricos, embora tal pareça estar relacionado com o RMA, o contrato através do qual os bancos estabelecem entre si as relações e as moedas de referência, como o euro (EUR), o dólar norte-americano (USD) e o zloty polaco (PLN).

Não obstante estas limitações, esta abordagem permitiu não só identificar corretamente tópicos-chave previamente detetados, como também reconhecer novos padrões, nomeadamente o Tópico 5, associado a pagamentos realizados por autodébito e às suas particularidades, e o Tópico 10, relacionado com erros na conta do beneficiário final.

A análise das métricas extrínsecas revela uma coerência de 0,50, valor que, embora ainda denote limitações na consistência semântica dos tópicos, representa uma ligeira melhoria face ao modelo de oito tópicos. Este resultado sugere que o aumento do número de tópicos tenha permitido uma segmentação temática um pouco mais precisa, ainda que persistam casos de redundância e sobreposição entre grupos. No conjunto, os resultados evidenciam que, apesar de o ganho em coerência ser modesto, a abordagem de dez tópicos contribuiu para uma identificação mais detalhada de incidentes, apenas com um impacto marginal na clareza interpretativa.

4.1.5 Considerações

A análise comparativa das diferentes configurações de modelação temática permitiu avaliar o desempenho dos modelos tanto sob uma perspetiva extrínseca, através da métrica de coerência, como sob uma perspetiva intrínseca, centrada na interpretabilidade e relevância dos tópicos identificados. De uma forma geral, as soluções com três e seis tópicos revelaram-se as mais equilibradas, apresentando bons resultados de coerência e uma interpretação consistente dos temas. O modelo de três tópicos destacou-se pela separação clara entre grupos, enquanto o modelo de seis tópicos, apesar de alguma sobreposição, ofereceu maior diversidade temática e uma compreensão mais detalhada das causas dos incidentes. Assim, considera-se que a configuração de seis tópicos constitui a solução mais adequada para este estudo, ao conjugar granularidade analítica com clareza interpretativa.

As abordagens com oito e dez tópicos, embora tenham registado valores inferiores de coerência e apresentado maior sobreposição temática, mostraram-se úteis na identificação de causas menos frequentes ou incidentes mais específicos que não tinham sido detetados pelas soluções mais agregadas. Ainda assim, a redundância observada confirma que o número de tópicos gerado é excessivo face à dimensão do conjunto de dados, conduzindo à repetição de padrões e a uma menor consistência semântica.

Em síntese, a utilização conjunta dos diferentes modelos permite uma visão mais abrangente do fenómeno em análise. O modelo de seis tópicos destaca-se como o mais equilibrado e esclarecedor, enquanto os modelos de maior granularidade, quando utilizados de forma complementar, acrescentam valor, ao revelarem padrões adicionais que aprofundam o entendimento global dos incidentes.

4.2 BERTopic

O BERTopic é um modelo de geração de tópicos que combina técnicas de *transformers* com métodos de redução de dimensionalidade e agrupamento, permitindo identificar de forma mais consistente estruturas semânticas em conjuntos de textos. Proposto por Grootendorst (2022), o BERTopic baseia-se na utilização de representações vetoriais de texto, também conhecidas como *embeddings*, geradas por modelos pré-treinados como o *Bidirectional Encoder Representations from Transformers* (BERT). Estas representações captam o significado contextual das palavras e frases, possibilitando a deteção de relações semânticas mais subtis do que obtidas por modelos puramente estatísticos, como o LDA. Assim, o BERTopic não depende da frequência de co-ocorrência de termos, mas sim da proximidade semântica entre as representações vetoriais, o que lhe confere uma capacidade acrescida de identificar tópicos com maior coerência semântica.

O funcionamento do modelo envolve várias etapas. Em primeiro lugar, o texto é transformado em *embeddings* utilizando modelos de linguagem baseados em *transformers*. De seguida, é aplicada uma técnica de redução de dimensionalidade para preservar a estrutura semântica do espaço vetorial num número reduzido de dimensões. Posteriormente, os vetores são agrupados por meio do algoritmo HDBSCAN (*Hierarchical Density-Based Spatial Clustering of Applications with Noise*), que permite identificar agrupamentos densos de documentos semanticamente semelhantes, ignorando *outliers*. Por fim, o modelo gera as palavras mais representativas de cada grupo, recorrendo à métrica *class-based Term Frequency–Inverse Document Frequency* (c-TF-IDF), que quantifica a relevância dos termos dentro de cada tópico.

Importa referir que o *BERTopic* é desenvolvido de forma modular, permitindo a substituição e personalização de cada um dos seus componentes principais, *embeddings*, método de redução de dimensionalidade, algoritmo de *clustering* e técnica de extração de palavras. Esta flexibilidade torna o modelo altamente configurável, possibilitando a adaptação a diferentes tipos de texto, volumes de dados e objetivos analíticos, o que o torna especialmente útil em contextos de investigação aplicada, como o deste projeto.

Numa fase inicial, o *BERTopic* foi implementado recorrendo ao *HDBSCAN* como método de agrupamento de tópicos, conforme sugerido pela revisão da literatura. A escolha deste algoritmo justifica-se pela sua capacidade de identificar automaticamente o número ideal de grupos, uma característica particularmente útil numa fase exploratória em que o número de tópicos latentes é desconhecido. Posteriormente, foram testadas soluções complementares utilizando o *K-Means* como método de *clustering*, permitindo controlar explicitamente o número de tópicos e avaliar diferentes configurações do modelo.

A presente implementação do *BERTopic* recorre ao UMAP como técnica de redução de dimensionalidade, uma vez que, tal como na literatura, esta abordagem preserva de forma eficiente as relações semânticas entre documentos no espaço vetorial, mesmo após a sua projeção em dimensões inferiores.

A avaliação dos modelos foi conduzida sob duas perspetivas complementares. Do ponto de vista extrínseco, recorreu-se à métrica de *coerência* (que possibilita uma comparação direta com os resultados obtidos através do modelo LDA), bem como à *diversidade* (que avalia a variabilidade lexical entre tópicos) e ainda à *perplexidade*, que mede a capacidade preditiva do modelo ao avaliar o como este generaliza para dados não observados, tendo em vista que valores mais baixos indicam um melhor ajustamento estatístico. Já a avaliação intrínseca assume um papel central, dado tratar-se de um modelo de natureza não supervisionada,

sendo baseada na análise qualitativa dos tópicos e no conhecimento de negócio do autor para determinar a sua relevância e interpretabilidade.

Dado que a metodologia adotada se baseia em *embeddings* e considerando a natureza dos dados, redigidos em inglês “Corporativo”, não pareceu ser necessário aplicar um processo de pré-tratamento textual. Uma vez que o modelo é capaz de captar as nuances semânticas e contextuais das palavras, a remoção de *stop words* não traria benefícios adicionais, podendo inclusive eliminar termos relevantes para a interpretação dos tópicos.

4.2.1 BERTopic utilizando HDBSCAN

No primeiro teste, o BERTopic foi implementado com a configuração genérica, alterando-se apenas o modelo de *embeddings*, substituindo o modelo base pelo FinBERT, mais adequado para textos de teor financeiro. Esta adaptação permitiu captar com maior precisão as nuances semânticas específicas do vocabulário técnico presente nos incidentes analisados. A aplicação deste modelo resultou na identificação de três tópicos principais, além de um tópico *outlier*, que agregou 18 documentos, correspondendo a aproximadamente 6,6% do total de registros considerados.

As métricas obtidas para esta configuração revelam uma coerência de 0,37, uma perplexidade de 1,19 e uma diversidade de tópicos de 0,73. Embora os valores de perplexidade e diversidade indiquem um desempenho estável e uma boa dispersão temática entre os tópicos, o valor de coerência situa-se muito aquém do esperado, registrando valores inferiores a todos os modelos LDA previamente testados. Esta diferença poderá estar relacionada com as características metodológicas distintas dos modelos, uma vez que o LDA se baseia em distribuições probabilísticas de palavras, enquanto o *BERTopic* utiliza representações vetoriais semânticas, que poderão afetar os resultados obtidos.

Para a visualização e interpretação dos tópicos, é possível identificar os documentos mais representativos de cada um, analisar as palavras principais associadas e respetiva relevância, observar a distribuição dos tópicos num mapa de análise de componentes principais e examinar as relações hierárquicas entre *clusters* resultantes do processo de agrupamento efetuado pelo HDBSCAN.

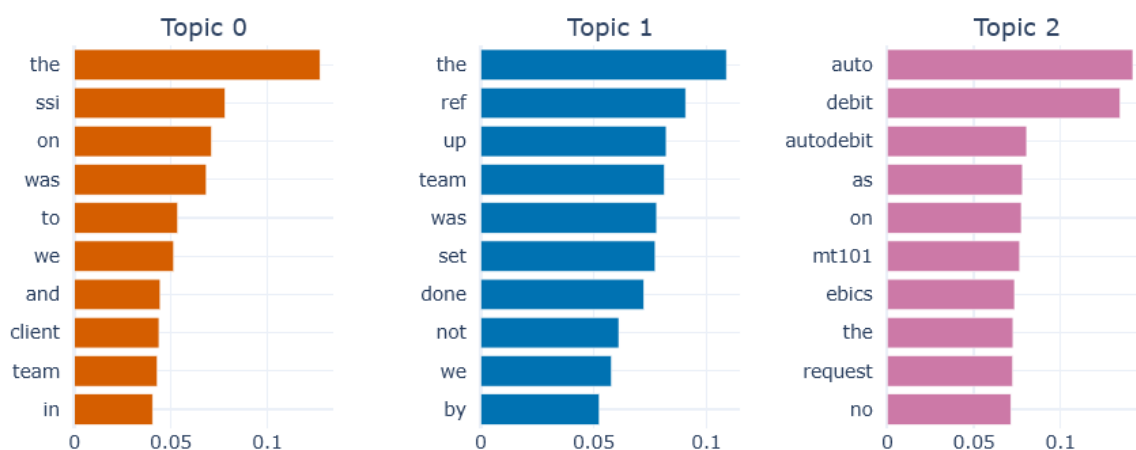


Figura 13: Termos mais relevantes por tópico no modelo BERTopic com 3 tópicos

A partir da análise da Figura 13, e recorrendo aos documentos representativos de cada tópico, foi possível classificar os temas identificados da seguinte forma:

- Tópico 0: Instruções do cliente não consideradas na atualização de SSI
- Tópico 1: *Setup* antigo, feito antes da equipa ter iniciado a atividade
- Tópico 2: Incidentes relacionados com autodebito

O primeiro tópico abrange 72,3% dos documentos, enquanto o segundo representa 15,1% e o terceiro apenas 5,9%, correspondendo a 16 documentos. Esta distribuição desequilibrada dos tópicos, aliada à presença de termos genéricos nos tópicos 0 e 1, evidencia a dificuldade do modelo em identificar causas concretas nos incidentes descritos.

O primeiro tópico representa 72,3% dos documentos, enquanto o segundo representa 15,1 e o terceiro representa 5,9% com apenas 16 documentos associados a este tópico. A distribuição irregular dos tópicos, em conjunto com os termos genéricos encontrados nos tópicos 0 e 1 mostram que estes modelos têm dificuldade em encontrar causas concretas dentro dos documentos.

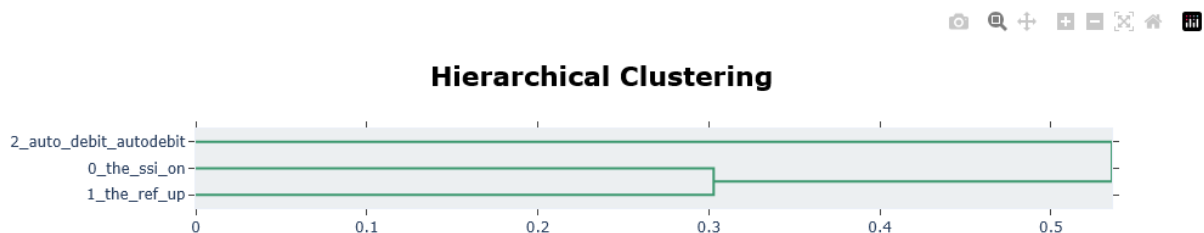


Figura 14: Estrutura hierárquica dos clusters de tópicos gerados pelo modelo BERTopic com 3 tópicos

A Figura 14 mostra a formação hierárquica dos *clusters*, onde se observa que os tópicos 0 e 1 se agrupam no mesmo ramo da árvore, indicando proximidade semântica entre si. A formação hierárquica baseia-se na distância derivada da similaridade do cosseno, onde valores mais próximos de 1 indicam maior proximidade semântica entre tópicos. A similaridade do *cosseno*, utilizada como medida de afinidade entre vetores, é definida matematicamente por:

$$\text{similaridade do cosseno} = \frac{\vec{A} \cdot \vec{B}}{\|\vec{A}\| \|\vec{B}\|}$$

A distância usada no dendrograma corresponde a:

$$\text{distância} = 1 - \text{similaridade do cosseno}$$

Deste modo:

- valores inferiores a 0,2 indicam tópicos praticamente idênticos
- valores entre 0,2 e 0,4 representam tópicos distintos, mas ainda dentro da mesma esfera temática
- valores entre 0,4 e 0,7 traduzem uma relação fraca, sugerindo apenas possíveis categorias amplas

Tal relação demonstra que os dois tópicos mais genéricos, que representam 87,4% dos documentos analisados, apresentam reduzida capacidade de diferenciação temática, limitando a eficácia do modelo na identificação de padrões relevantes.

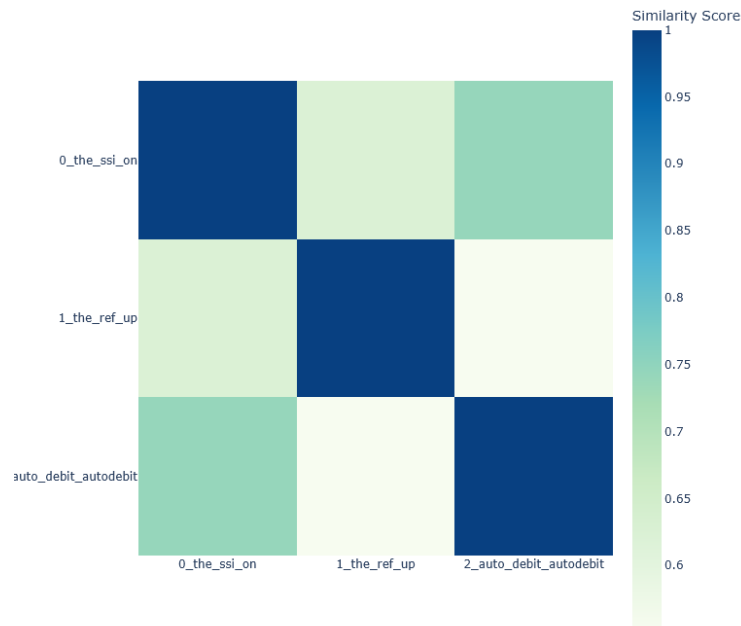


Figura 15: Mapa de Similaridade entre tópicos no modelo BERTopic com 3 tópicos

Por sua vez, a Figura 15 demonstra que não existem correlações significativas entre os tópicos, sugerindo uma separação temática adequada entre os grupos identificados.

Após esta primeira execução, concluiu-se que os resultados obtidos eram excessivamente genéricos, com um dos tópicos a abranger a maior parte do conjunto de dados, o que limitou a capacidade do modelo em distinguir causas específicas dos incidentes. Para ultrapassar esta limitação, procedeu-se a ajustes nos parâmetros do UMAP e do HDBSCAN, com o objetivo de tornar o *BERTopic* mais sensível a variações semânticas e mais adequado a conjuntos de dados pequenos e heterogêneos, como é o caso dos relatórios de incidentes operacionais.

A generalização excessiva observada nos resultados iniciais poderá estar relacionada com uma redução de dimensionalidade demasiado agressiva dos dados, que pode ter comprimido excessivamente o espaço vetorial, dificultando a distinção entre temas semanticamente próximos. Adicionalmente, a configuração genérica do HDBSCAN poderá ter contribuído para este efeito, ao impor um tamanho mínimo de documentos por *cluster*, o que pode ter levado à fusão de tópicos diferentes num mesmo grupo.

As alterações introduzidas no UMAP visaram melhorar a preservação da estrutura semântica durante a redução de dimensionalidade. O parâmetro (*n_neighbors*) relativo ao número de vizinhos foi aumentado para 30, de modo a captar relações globais entre os documentos, evitando que *clusters* semelhantes fossem separados de forma artificial. O número de dimensões projetadas foi definido como 5, permitindo reter mais informação contextual dos *embeddings*. Além disso, o parâmetro (*min_dist*) relativo à distância mínima foi ajustado para 0,1, promovendo maior coesão interna nos clusters, e manteve-se a métrica do tipo *cosine*, mais adequada para medir similaridade entre representações textuais.

Relativamente ao HDBSCAN, os parâmetros foram ajustados para equilibrar a formação de *clusters* mais pequenos e específicos com a redução de ruído. O tamanho mínimo de *cluster* foi definido como 3, permitindo identificar tópicos com menor representatividade, mas potencialmente relevantes. O parâmetro (*min_samples*) relativo ao número mínimo de *samples* foi aumentado para 5, tornando o modelo mais restritivo na criação de grupos e reduzindo a probabilidade de *clusters* espúrios. Estas alterações resultaram num modelo mais

refinado e interpretável, capaz de captar padrões temáticos mais granulares sem comprometer a coerência semântica geral.

Foram realizados testes aplicando as otimizações do HDBSCAN e do UMAP, tanto em conjunto como de forma independente, de modo a avaliar o impacto de cada configuração no desempenho do modelo.

Apesar de o *BERTopic*, pela sua natureza baseada em *embeddings*, não exigir um pré-processamento extensivo do texto, a análise inicial dos resultados obtidos com as configurações testadas revelou a presença recorrente de meses e de certos números entre os termos mais representativos de vários tópicos. Estes elementos não acrescentavam informação relevante para a interpretação semântica dos incidentes e, por esse motivo, foram removidos do vocabulário utilizado pelo modelo. Ainda assim, foram introduzidas exceções relativamente aos números excluídos, de modo a preservar aqueles que possuem significado operacional no contexto dos incidentes analisados, nomeadamente 50, 57, 58 e 59, que correspondem a campos específicos das mensagens SWIFT, bem como 101, 103 e 202, que identificam tipos de mensagem dentro do mesmo sistema.

A Tabela 7 apresenta os resultados das métricas obtidas em ambas as situações.

Tabela 7: Avaliação de métricas extrínsecas entre os modelos BERTopic

Métrica	<i>HDBSCAN + UMAP</i>	<i>HDBSCAN</i>	<i>UMAP</i>	<i>BERTopic base</i>
Nº de Tópicos	17	18	2	3
Coerência	0,4445	0,5353	0,4135	0,37
Diversidade	0,6955	0,9750	0,8500	1,19
Perplexidade	1,5488	1,3999	1,1294	0,73

O modelo configurado com UMAP apresentou o desempenho mais baixo entre as três abordagens, registando apenas valores de coerência e perplexidade ligeiramente superiores aos obtidos pelo *BERTopic base*. Além disso, o modelo produziu apenas dois tópicos principais, ambos com temáticas semelhantes às identificadas no *BERTopic base*, nomeadamente incidentes relacionados com autodébito, que representaram 14,76% dos documentos analisados. O tópico predominante, de natureza genérica, abrangeu 82,66% do conjunto de dados, revelando baixa capacidade discriminativa do modelo em identificar causas específicas dos incidentes.

Entre a abordagem que recorre apenas ao HDBSCAN e a que combina HDBSCAN com UMAP, a configuração que utiliza exclusivamente o HDBSCAN apresentou um aumento significativo nos valores de coerência e diversidade, o que indica tópicos mais consistentes, semanticamente mais claros e com maior variabilidade lexical entre si.

Em contrapartida, a integração do UMAP resultou num aumento significativo da perplexidade, sugerindo uma maior dificuldade do modelo em generalizar para novos dados e uma menor estabilidade na identificação dos padrões semânticos subjacentes.

Optou-se, assim, por adotar a abordagem que recorre apenas às otimizações do HDBSCAN, uma vez que este apresentou resultados superiores nas métricas extrínsecas e também por apresentar um tópico adicional que poderá ser relevante na identificação de mais uma causa associada a incidentes operacionais.

Tabela 8: Frequência e Proporção de Documentos por Tópico

Métrica	Outlier	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17
Proporção (%)	16	12,4	11,3	8,7	5,1	5,1	5,1	4,7	4	3,6	2,9	2,9	2,9	2,5	2,5	2,2	2,2	2,2	2,2

Conforme a Tabela 8, o modelo identificou dezoito tópicos, que no seu conjunto representam cerca de 84% dos incidentes analisados, enquanto os restantes foram classificados como *outliers*. A distribuição dos incidentes pelos diferentes tópicos varia entre 12,4% e 2,2%, sendo que os três tópicos mais representativos concentram aproximadamente 32,5% dos incidentes identificados pelo sistema.

Esta variação permite reconhecer simultaneamente os temas mais recorrentes, associados a incidentes com maior expressão, e os temas menos frequentes, que apesar de representarem uma proporção reduzida do total, podem ainda assim ser operacionalmente relevantes, uma vez que correspondem a ocorrências que facilmente poderiam passar despercebidas na análise manual.

O grupo classificado como *outlier* representa 16% dos casos, correspondendo a cerca de um quinto das observações. Este valor deve ser interpretado com cautela, pois um volume elevado de *outliers* pode influenciar a perceção da relevância relativa dos tópicos identificados. Contudo, importa reconhecer que, em modelos de *topic modelling*, a presença de documentos que não se ajustem claramente a nenhum tópico é inevitável, quer devido às ambiguidades e subtilezas linguísticas, quer pela existência de descrições pouco específicas ou com informação insuficiente para permitir a sua integração coerente num *cluster* semântico. Assim, o número elevado de *outliers* poderá refletir limitações na qualidade ou fiabilidade dos comentários analisados, do que um problema intrínseco do modelo semântico.

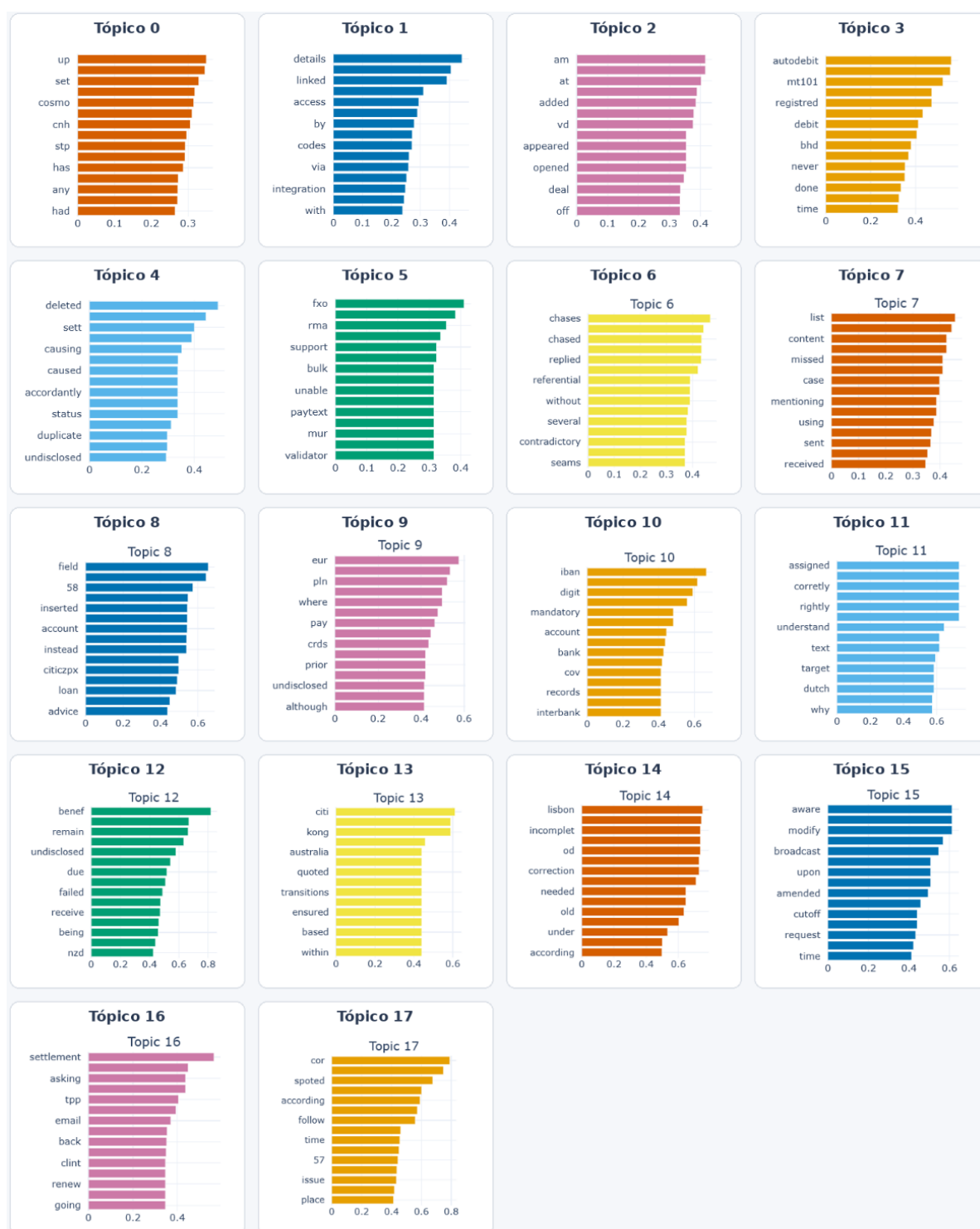


Figura 16: Termos mais relevantes por tópicos no modelo BERTopic com 18 tópicos

Através da análise da Figura 16, em conjunto com a análise dos 3 documentos representativos atribuídos pelo modelo a cada tópico, foi possível identificar os seguintes tópicos:

- Tópico 0: Notificação de *Alert* - *HKD/CNH*
- Tópico 1: *Alert details* incorretos
- Tópico 2: *SSI_NONE* - *booking* tardio
- Tópico 3: *Autodebit* MT101 – Registo/Aplicação Incorreta
- Tópico 4: *Undisclosed Client*: Nome apagado do sistema
- Tópico 5: Alterações Não Aplicadas (MUR ccy - Bulk FXO)
- Tópico 6: Resposta tardia do cliente
- Tópico 7: Instruções do Cliente Não Consideradas
- Tópico 8: Erros nos Campos 57/58 (Interpretação Incorreta)
- Tópico 9: *Undisclosed* – SSI Errada / BENF Desatualizado (EUR–PLN)
- Tópico 10: IBAN Inválido (Dígito em Falta)
- Tópico 11: Incidentes Atribuídos Indevidamente
- Tópico 12: *Undisclosed* – Nome Beneficiário Não Atualizado
- Tópico 13: Migração: BIC no Campo Errado - Falha de Comunicação
- Tópico 14: *Old setup*: Setup feito antes da equipa tomar responsabilidade
- Tópico 15: *Broadcast Update*: Atualização Não Comunicada / Comunicação Tardia
- Tópico 16: Falha na comunicação inter-equipas
- Tópico 17: *Setups* em conformidade

O modelo demonstrou capacidade para identificar de forma eficaz diversas causas associadas aos incidentes de produção, gerando tópicos distintos e bem delimitados tematicamente. Foi igualmente capaz de distinguir corretamente dois incidentes relacionados com a temática dos *undisclosed clients*, sem ocorrência de sobreposição semântica entre os tópicos, conseguindo separar causas que partilhem vários dos mesmos termos frequentes. O tópico 4 refere-se a incidentes em que o nome *undisclosed* escolhido pelo cliente foi eliminado do sistema, originando uma falha por omissão de informação, enquanto o tópico 12 descreve situações em que o nome não foi devidamente atualizado, resultando em pagamentos processados com dados desatualizados.

O modelo identificou igualmente tópicos relacionados com atrasos em pagamentos e incidentes cuja responsabilidade não pertence à equipa que reportou o caso. Embora tenham sido removidos previamente todos os incidentes sem responsabilidade direta da equipa e excluídos também os casos em que a causa estivesse associada a atrasos, alguns registos acabaram por ser incluídos devido a classificações incorretas. Os tópicos 6, 11, 14 e 17 exemplificam situações em que a origem do incidente não estava relacionada com a atuação da equipa. Esta ocorrência evidencia uma das limitações inerentes à análise de dados auto-reportados, sobretudo quando não existe uma orientação clara sobre o preenchimento adequado da base de dados. Ainda assim, estes três tópicos representam apenas uma pequena porção dos incidentes analisados, o que sugere que, de forma geral, o impacto dos dados incorretamente reportados é reduzido.

Pode observar-se alguma sobreposição temática entre certos tópicos, sobretudo nos que dizem respeito à atribuição incorreta de responsabilidade à equipa. Esta sobreposição poderá resultar não só da reduzida dimensão de alguns tópicos, mas também do facto de vários incidentes terem sido indevidamente associados à equipa quando na realidade estavam relacionados com limitações do sistema ou com erros do próprio cliente. Assim, embora o aumento do número de tópicos permita identificar causas adicionais, acaba também por fragmentar temas que, do ponto de vista da análise humana, não apresentam uma diversidade suficientemente significativa para justificar essa separação.

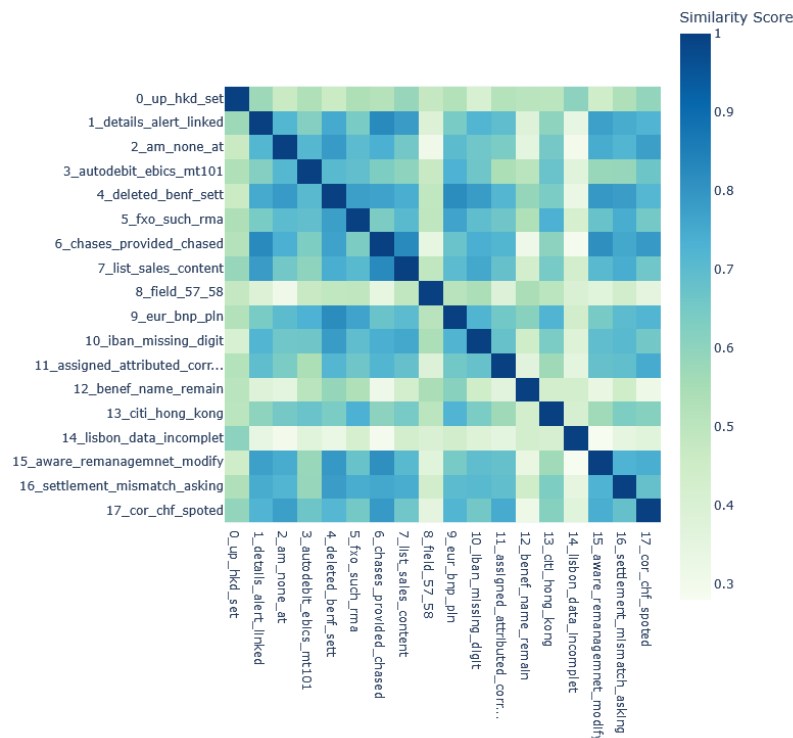


Figura 17: Mapa de Calor da Similaridade entre Tópicos

Conforme ilustrado na Figura 17, não se observam correlações fortes entre os tópicos, o que sugere que o modelo conseguiu distinguir adequadamente os diferentes temas. Esta fraca similaridade global indica que os tópicos são semanticamente bem separados, reduzindo o risco de sobreposição e reforçando a qualidade da segmentação temática obtida.

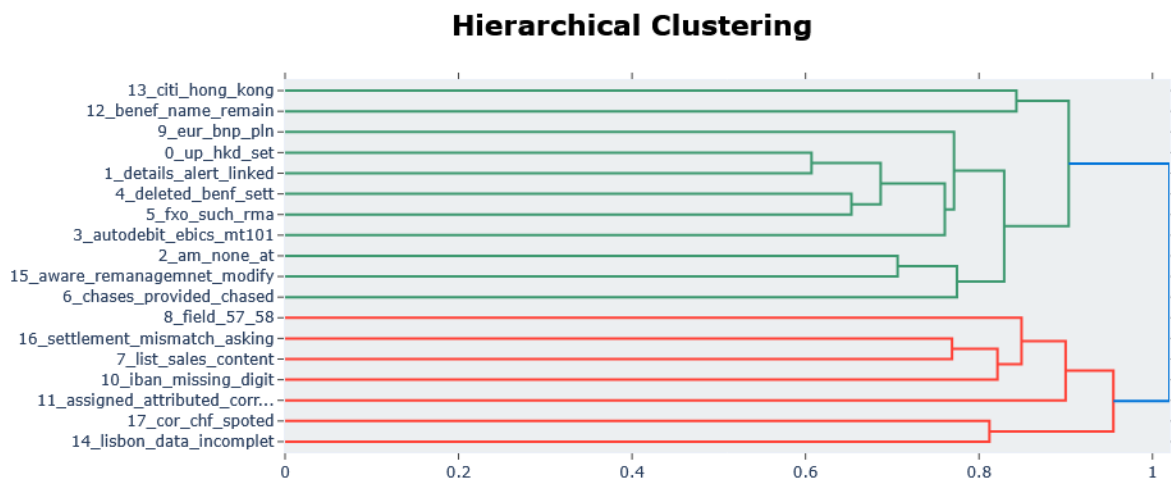


Figura 18: Dendrograma de Agrupamento Hierárquico dos Tópicos

Através da Figura 18 é possível analisar a formação hierárquica dos *clusters*, permitindo identificar possíveis similaridades entre tópicos que, em teoria, poderiam ser agrupados em super-tópicos.

Verifica-se que todas as ligações entre tópicos apresentam valores superiores a 0,6, o que indica que os tópicos se encontram bem separados do ponto de vista semântico. Esta interpretação é reforçada pela matriz de similaridade representada na Figura 17. Face a estes resultados, não existem indícios que justifiquem a agregação dos tópicos em super-tópicos.

Ainda assim, caso se pretendesse analisar apenas as relações mais fortes disponíveis, entre 0,6 e 0,7, poderiam ser identificadas algumas categorias de nível mais elevado, embora bastante amplas.

Estes valores permitiriam identificar duas possíveis agregações de tópicos. A primeira corresponde aos tópicos 0, 1, 4 e 5, podendo ser subdividida em dois conjuntos: o primeiro integra os tópicos 0 e 1, enquanto o segundo reúne os tópicos 4 e 5. A segunda agregação é composta pelos tópicos 2 e 15.

O primeiro grupo pode ser enquadrado como ‘Falhas sistémicas ou de processos’, uma vez que reúne tópicos cuja origem está associada ao funcionamento dos sistemas ou à execução dos processos operacionais. Seguindo a subdivisão referida anteriormente, o primeiro subgrupo pode ser descrito como ‘Incidentes ALERT’, dado que ambos os tópicos se relacionam diretamente com o sistema de atualizações de instruções *Alert*. O segundo subgrupo pode ser caracterizado como ‘Falhas em processos automáticos’, por envolver situações em que ocorreram pagamentos com informação incompleta ou desatualizada.

Por sua vez, o segundo grupo está claramente associado a situações decorrentes de uma ação tardia por parte da equipa, motivada por um desfasamento temporal entre o momento da necessidade e a realização das tarefas, seja por *bookings* efetuados próximo do *cut-off*, seja por notificações de *broadcast* emitidas demasiado tarde. Assim, este conjunto pode ser designado como ‘notificação tardia do pedido’.

O modelo apresentou o desempenho mais consistente entre todas as abordagens testadas, tanto probabilísticas como semânticas. Este resultado observa-se quer na avaliação extrínseca, através de métricas como coerência, diversidade e perplexidade, quer na avaliação intrínseca, pela facilidade de interpretação dos tópicos e pela reduzida sobreposição temática. O modelo demonstrou ainda capacidade para gerar tópicos com termos dominantes semelhantes, mas semanticamente distintos, mostrando que os métodos baseados em *embeddings* conseguem captar nuances linguísticas que permitiriam distinguir documentos que, num modelo probabilístico, seriam agrupados no mesmo tópico, como é o caso dos incidentes associados à plataforma ALERT ou a *undisclosed clients*.

Apesar do bom desempenho do modelo, subsistiram algumas limitações. O modelo identificou incidentes relacionados com atrasos em pagamentos e situações em que a responsabilidade não era diretamente atribuível à equipa analisada. Esta limitação decorre da própria natureza dos dados autoreportados: mesmo após o tratamento e filtragem das variáveis, alguns incidentes foram incorretamente categorizados na origem e acabaram por ser incluídos na análise. Assim, recomenda-se, como medida de aperfeiçoamento futuro, a definição prévia de regras mais claras e padronizadas para o preenchimento desta informação, de modo a reforçar a qualidade dos dados.

Ainda assim, o modelo conseguiu identificar um conjunto relevante de tópicos alinhados com as principais causas de incidentes na organização, facilmente interpretáveis através dos termos e documentos dominantes. Isto permitiu aprofundar a compreensão das fontes dos incidentes e criar uma base mais sólida para discussões sobre mitigação e prevenção.

Atendendo à qualidade dos resultados obtidos e ao potencial de interpretação associado, este modelo foi selecionado para a análise complementar com apoio de um *Large Language Model*, com o objetivo de reforçar o entendimento das informações extraídas e mitigar possíveis enviesamentos de interpretação individual.

4.2.2 Análise complementar com apoio de LLM

Uma das principais dificuldades em se trabalhar com modelos de tópicos não supervisionados é a avaliação e classificação correta. Não existindo categorias pré-definidas, a classificação dos tópicos depende fundamentalmente da interpretação pessoal, exigindo um forte conhecimento de negócio para uma boa classificação dos incidentes. Ainda assim, esta avaliação não deixa de estar sujeita a um viés pessoal que pode afetar a forma como os tópicos são classificados.

Por esse motivo, a análise complementar com LLM é um passo recomendado pela documentação oficial do *BERTopic* (Grootendorst, 2024) com o objetivo reduzir o impacto do viés pessoal. Este passo permite reforçar a análise individual em momentos onde a análise do modelo de linguagem se encontra em convergência com a análise individual do autor. Permite também identificar novos padrões que podem passar despercebidos à primeira análise.

Foi utilizado o modelo de linguagem da *OpenAI*, *GPT 5.1*, capaz de ler e interpretar um grande volume de dados. Numa fase inicial, introduziu-se o contexto da análise, bem como o contexto operacional, fazendo a apresentação de conceitos específicos das realidades operacionais. Depois, apresentou-se a Figura 16, com os termos dominantes associados a cada tópico, bem como os 3 documentos representativos associados a cada tópico, dado que estes foram os elementos utilizados para se interpretar e nomear os tópicos pela análise do autor, e pediu-se ao modelo que fizesse o mesmo.

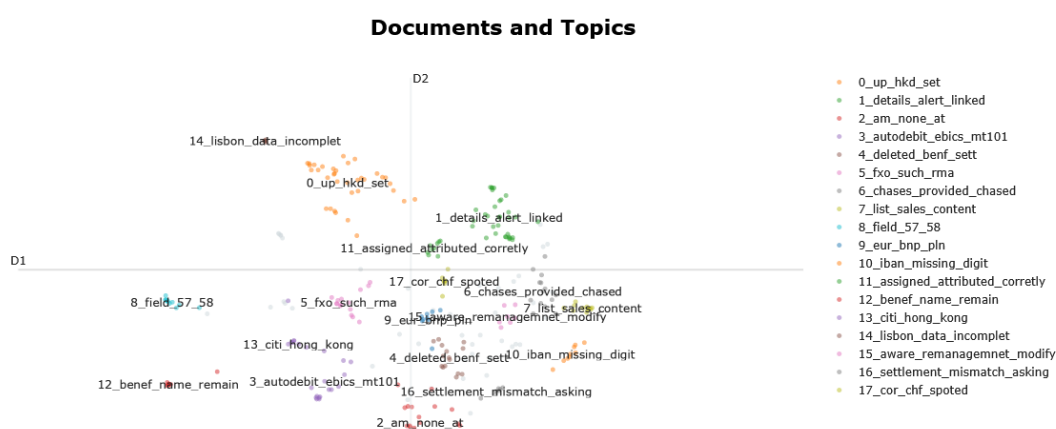


Figura 19: Projeção Bidimensional dos Documentos e Distribuição dos Tópicos

Através dos resultados gerados pelo modelo, fez-se uma análise comparativa que permitiu entender pontos de convergência e divergência entre os resultados do modelo e a análise pessoal.

Por fim, pediu-se ao modelo que interpretasse a Figura 19, com o objetivo de encontrar possíveis grupos de tópicos ou novas formas de entender a natureza dos incidentes, tentando entender e explicar as dimensões criadas.

Tabela 9: Avaliação de resultados feita pelo GPT5.1

Tópico	Nome proposto	Categoria principal
0	Erros de Processo Interno / Coordenação	Processo
1	Ligação/Atualização Incorreta no <i>Alert/Cosmo</i>	Sistemas externos
2	<i>Late Booking</i> por <i>SSI_NONE</i> / <i>Cut-off</i>	<i>Cut-off</i>
3	<i>Setup</i> Incorreto de MT101/ <i>Autodebit</i>	<i>Autodebit</i>
4	<i>Setup</i> Duplicado/Apagado/Inválido	Qualidade de dados
5	<i>Updates</i> Pendentes Não Aplicados	Processo / Governança
6	Cliente responde tarde / após COT	Cliente
7	Informação do Cliente Ignorada	Cliente
8	Erro de Interpretação SWIFT	Qualidade de dados
9	<i>Setup</i> incompleto / <i>missing BENF</i>	<i>Undisclosed</i>
10	IBAN inválido / problema técnico	Dados técnicos
11	Incidente não pertencente ao REF	Governança
12	Falta de Atualização do = <i>BENF</i> em <i>Undisclosed</i>	<i>Undisclosed</i>
13	BIC incorreto por informação errada	<i>Custodian</i>
14	<i>Setup</i> herdado / dados incompletos	Dados herdados
15	Mudança de SSI não comunicada	<i>Custodian</i>
16	Mau entendimento da requisição	Comunicação
17	<i>Setup</i> correto / sem erro	Não incidente

O modelo apresentou os resultados em formato de tabela, nomeando cada tópico e atribuindo uma categoria principal a cada um dos incidentes, conforme podemos ver na tabela 9. Comparando os nomes atribuídos pelo modelo com os nomes sugeridos pela análise individual, podemos notar que existe uma convergência em praticamente todos os tópicos com a exceção de 0, 4 e 16.

No caso do tópico 0, a diferença de classificação explica-se pelo facto de a análise individual ter atribuído mais importância aos termos dominantes atribuídos ao tópico, enquanto LLM atribuiu maior importância aos documentos representativos, concluindo de forma correta que os 3 relatórios de incidente abordam falhas no processo, seja por incapacidade de lidar com o volume, como por dificuldades de coordenação com outras equipas e até com o cliente.

O contrário acontece no tópico 4, em que o modelo ignora que todos os incidentes representativos têm o termo '*undisclosed*' e tenta explicar a ação que levou à falha operacional com base nos termos dominantes. Podemos admitir que os nomes atribuídos permitiram uma visão mais completa destes incidentes, em que LLM se foca mais em explicar as ações que levaram ao incidente, enquanto a análise individual reforça a categoria de incidentes relacionada com o campo do nome do cliente.

Por fim, relativamente ao tópico 16, apesar de os nomes atribuídos serem relativamente semelhantes, o LLM deu maior relevância aos termos dominantes, focando-se na requisição, um termo relativo à renovação de um acordo para um pagamento *TPP*, enquanto a análise individual atribui mais relevância à análise dos documentos representativos, que demonstram falhas na comunicação entre equipas. Os documentos representativos representam metade deste tópico, motivo pelo qual, ambas as interpretações são complementares.

O modelo de linguagem analisou também possíveis agrupamentos entre tópicos. Embora tenha identificado certos tópicos com elementos comuns que, à partida, poderiam ser

agregados, verificou-se que cada um possuía sempre um elemento semântico distintivo na causa do incidente. Um exemplo ilustrativo surge nos tópicos associados a *undisclosed clients*: no primeiro caso, o incidente resulta da remoção incorreta do indicador de cliente não divulgado, enquanto no segundo o problema decorre da alteração do nome do beneficiário, que não foi atualizada pela equipa. Assim, apesar das semelhanças superficiais, as causas subjacentes diferenciam-se de forma suficientemente clara para justificarem a manutenção dos tópicos como unidades independentes.

De seguida, foi pedido ao modelo de LLM que analisasse a Figura 19, de modo que fizesse uma análise sobre a posição dos tópicos no mapa bidimensional, de modo a tentar entender o que determina as dimensões e como estas influenciam o tipo de incidentes identificados. Para auxiliar no processo de interpretação, foi explicada no *prompt* a posição dos tópicos dentro das dimensões e dentro de cada quartil gerado pelo gráfico.

Para a dimensão da esquerda para a direita, o sistema sugeriu como nome ‘Eixo Dados Técnicos x Ação Operacional’, justificando que os incidentes mais à esquerda estariam relacionados com erros estruturais, dados incorretos ou *setups* antigos, enquanto os incidentes à direita estariam relacionados com falhas humanas, atrasos e má interpretação de instruções.

Já a segunda dimensão é definida como ‘Eixo Causa Sistémica x Causa Operacional Direta’, demonstrando que os tópicos que estão mais em cima estão relacionados com incidentes estruturais e relacionados com os sistemas, enquanto que os que estão mais em baixo estão incidentes.

Esta análise realizada pelo modelo de linguagem *GPT 5.1* acaba por demonstrar algumas inconsistências em ambas as dimensões, com vários tópicos que não seguem a lógica proposta pelas mesmas. Isto deve-se à falta de contexto sobre a natureza dos incidentes, o que o leva a fazer suposições incorretas sobre a natureza dos mesmos. Esta situação ocorre mesmo providenciando o máximo de contexto relativo possível. A dificuldade do modelo em encontrar um padrão para as dimensões pode também ser explicada pela análise prévia das Figuras 17 e 18 que evidenciam separação temática entre os tópicos, seja pela fraca similaridade entre si, seja por a separação hierárquica dos tópicos ter valores de similaridade negativa muito elevados.

4.2.3 Considerações

Através das experiências realizadas conseguimos perceber que o BERTopic demonstrou resultados superiores aos dos testes realizados com modelos probabilísticos, a nível intrínseco, com tópicos mais fáceis de interpretar e mais bem separados tematicamente, evidenciando que a capacidade de identificar estruturas semânticas contribui para uma melhor interpretação dos resultados gerados.

Entre as experiências realizadas, a que se destacou com melhores resultados foi a que otimizou os parâmetros do HDBSCAN para permitir um número maior de tópicos com menos representatividade sem comprometer a coerência semântica. Para além destas alterações, foram removidos todos os meses e datas e reduziu-se o peso de *stop words*. Esta experiência obteve o maior valor de coerência registado com 0,53, obtendo 17 tópicos bem delimitados e separados tematicamente, permitindo descobrir causas relevantes dos incidentes que podem acabar por passar despercebidas por terem menos ocorrências.

Por fim, foi realizado um teste com um modelo LLM, onde se pediu ao modelo GPT5.1 da Open AI que interpretasse os resultados obtidos, de forma a reduzir o viés de interpretação causado pela análise individual. Através desta análise verificou-se que a grande maioria dos

nomes atribuídos pelo LLM corresponderam aos mesmos que tinham sido atribuídos na análise individual e ainda que, nos tópicos onde houve uma divergência, foi possível aprofundar o entendimento sobre os mesmos.

Apesar de a interpretação gerada pelo LLM se ter mostrado bastante relevante, sendo uma ferramenta útil para reforçar a análise e reduzir ambiguidades causadas por vieses, estes modelos apresentam limitações nos resultados, estando sempre dependentes do contexto fornecido no *prompt* utilizado. Por esse motivo, este passo deve ser visto como complementar e não substituto da análise individual substanciada com conhecimento de causa.

5 Validação dos Resultados através de Inquérito às Equipas

Uma das principais limitações de se trabalhar com modelos não supervisionados de extração de tópicos é a subjetividade dos resultados obtidos, bem como a sua aplicabilidade limitada. Apesar de existirem métricas extrínsecas que permitem avaliar de forma objetiva a coerência dos tópicos gerados, a atribuição de um rótulo é sempre uma análise pessoal e subjetiva e mesmo quando esta é fundamentada em conhecimento operacional, está sujeita a um desvio de interpretação. Além disso, a utilidade destes resultados para responder às questões de investigação e ao problema de negócio é também subjetiva na medida em que a utilidade percebida destes modelos depende de uma avaliação subjetiva dos resultados.

Por isso, revelou-se fundamental uma fase de validação dos resultados pelas equipas envolvidas nos incidentes de produção, permitindo responder às questões de investigação e ao problema de negócios, de modo a confirmar se estes modelos são de facto úteis para a discussão e possível prevenção das causas dos incidentes.

De forma a reduzir ao máximo vieses possíveis, foi inquirida toda a população operacional das equipas de *Referential* envolvidas nos incidentes de produção.

No âmbito desta análise, foram utilizados os tópicos identificados pelo modelo BERTopic, cuja configuração incluiu a otimização dos parâmetros do HDBSCAN.

5.1 Descrição do questionário e fundamentação das questões

O inquérito foi dividido em três secções distintas com o objetivo de responder a diferentes questões fundamentais para a validação dos resultados.

A primeira foi nomeada de ‘Avaliação Individual dos Tópicos’ e tem como objetivo avaliar a classificação feita do tópico e a compreensão temática do mesmo. Para isso, foram apresentados os termos dominantes de cada tópico presentes na Figura 16 e os incidentes representativos do tópico para responderem às questões

- ‘Com base na figura apresentada e nos incidentes representativos, considera clara a compreensão do tema abordado pelo tópico?’
- ‘O título atribuído ao tópico reflete corretamente o conteúdo dos incidentes?’

A primeira questão procura entender se o tópico é facilmente interpretável pelo inquirido ou se, pelo contrário, a análise feita para atribuir o nome do tópico é forçada. A segunda questão procura verificar se o nome atribuído faz sentido perante os resultados gerados.

Para esta secção foram apenas considerados os 7 tópicos com mais incidentes atribuídos, correspondendo a mais de metade dos incidentes identificados pelo modelo. Esta decisão foi feita para não tornar o inquérito demasiado exaustivo, comprometendo a qualidade de resposta dos inquiridos com demasiadas perguntas.

A segunda secção, ‘Avaliação Global dos Tópicos’, tem como objetivo avaliar os resultados do modelo de forma agregada, nomeadamente no que respeita à interpretação dos tópicos, à sua utilidade para a compreensão dos incidentes e ao seu potencial contributo para a prevenção e melhoria de processos.

Para aferir a qualidade da interpretação geral dos tópicos, foram incluídas as questões:

- ‘Observou tópicos que parecem demasiado semelhantes entre si?’
- ‘Identificou algum tópico cujo conteúdo seja demasiado diverso para estar agrupado?’

Sempre que o inquirido responda afirmativamente a qualquer uma destas perguntas, surge automaticamente a questão adicional ‘Se sim, quais?’, de forma a identificar concretamente quais os tópicos considerados problemáticos.

Para entender se a segmentação dos incidentes por tópicos melhora a compreensão dos incidentes operacionais e a das respetivas causas de risco foram feitas quatro questões.

A primeira pergunta, ‘A distribuição dos incidentes por tópico é útil para a compreensão das principais causas dos incidentes?’, procura entender se a organização dos incidentes por tópicos aumenta o conhecimento, de forma geral, sobre as causas de incidentes.

Já a segunda questão, ‘Considera que a segmentação por tópicos ajuda a entender as causas dos incidentes como um todo?’ com o objetivo entender se os resultados desenvolvidos pelo modelo reforçam, de forma abrangente, o conhecimento sobre as causas dos incidentes dentro da equipa

A terceira pergunta, ‘Os tópicos mais frequentes correspondem, na sua experiência, às áreas mais críticas?’, pretende avaliar em que medida a distribuição de frequência dos tópicos está alinhada com a perceção operacional sobre a frequência dos diferentes tipos de incidentes. Por fim, a questão ‘O modelo revelou alguma causa de risco que, na sua opinião, normalmente passaria despercebida?’ procura identificar a capacidade do modelo em identificar causas menos evidentes que poderiam passar despercebidas, onde caso o inquirido responda sim, é apresentada uma questão e em seguida uma segunda questão, perguntando quais os tópicos onde essas causas foram identificadas.

Quanto à utilidade dos modelos para prevenção e melhorias de processos foram levantadas 3 questões:

- ‘Na sua opinião, os tópicos ajudam a discutir ações preventivas dentro da equipa?’
- ‘Acredita que este tipo de análise ajudaria na identificação precoce de padrões de erro?’
- ‘Vê utilidade na repetição anual deste tipo de estudos com os dados atualizados?’

A última secção, denominada de ‘Percepção Geral’, procura responder à questão ‘Como avalia globalmente a qualidade da análise produzida?’ atribuindo-se uma notação de 0 a 5 na resposta. Para além desta avaliação, é incluída uma questão adicional, de resposta aberta, para que os participantes possam deixar comentários livres ou sugestões de melhoria relativamente ao modelo.

5.2 Análise e interpretação das respostas

Entre os sete tópicos analisados individualmente, a maioria dos inquiridos considerou que o título atribuído aos tópicos reflete corretamente o conteúdo dos incidentes (Figura 20), com a frequência das respostas positivas a variarem entre 78,6% e 57,1%; as respostas negativas a variarem entre 21,4% e 14,3% e por fim as respostas indecisas a variarem entre 7,1% e 14,3%. Entre estes, o tópico 6 destaca-se por ter o maior nível de concordância generalizada, com 78,6% da população a concordar com o título e apenas 2 inquiridos a discordarem, enquanto o tópico 2 se destaca pelo menor nível de concordância com apenas 57,1% a concordarem com o título atribuído e com a mesma percentagem de respostas negativas e indecisas (21,4%).

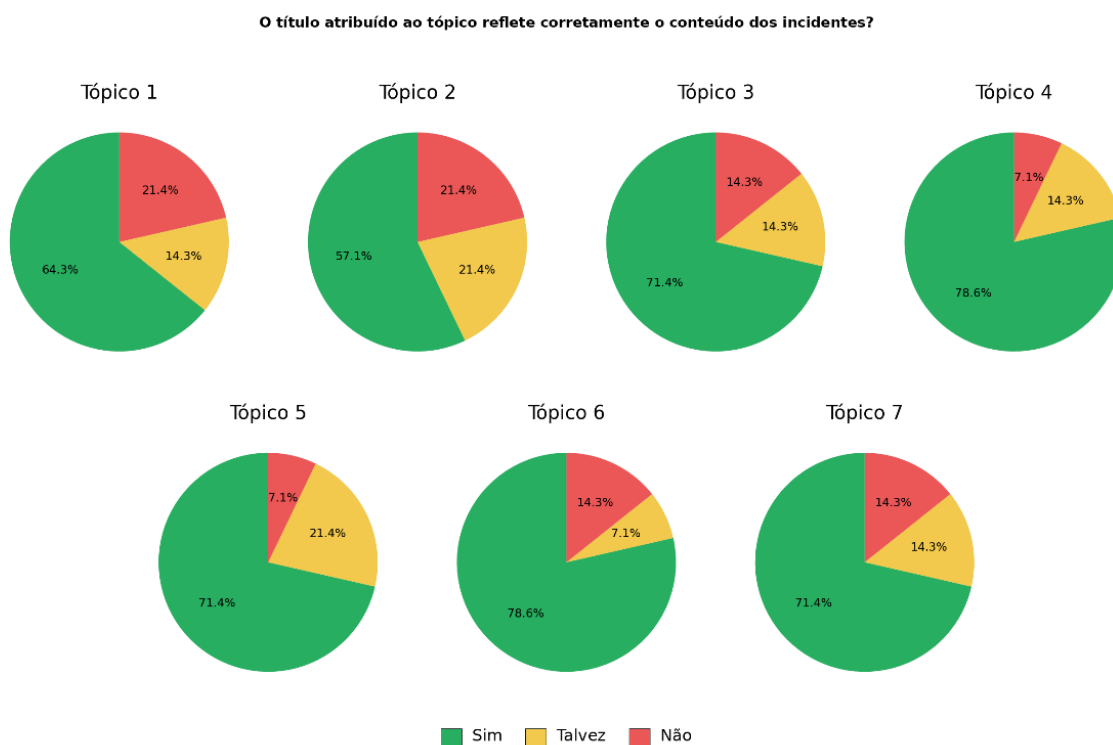


Figura 20: Avaliação da adequação do título atribuído a cada tópico face ao conteúdo dos incidentes

Relativamente à compreensão individual dos temas abordados pelos tópicos, conforme a Figura 21, os resultados indicam uma perceção globalmente positiva quanto à clareza dos tópicos apresentados, com a maioria das respostas concentrada nas categorias “Claro” e “Muito claro” em todos os tópicos avaliados, com respostas negativas “Pouco claro” e “Nada claro” em conjunto a variarem entre 28,4% e 14,2%. Esta distribuição sugere que, de um modo geral, os inquiridos conseguem compreender o tema central de cada tópico com base nos termos dominantes e nos incidentes representativos apresentados.

O tópico 7 destaca-se por reunir a maior proporção de respostas no nível mais elevado de clareza, mantendo simultaneamente uma baixa incidência de avaliações negativas. Por outro lado, o tópico 6 apresenta o desempenho mais fraco, com 28,5% dos resultados negativos, com 7,1% a considerarem-no como nada claro.

Com base na figura apresentada e nos incidentes representativos, considera clara a compreensão do tema abordado pelo tópico?

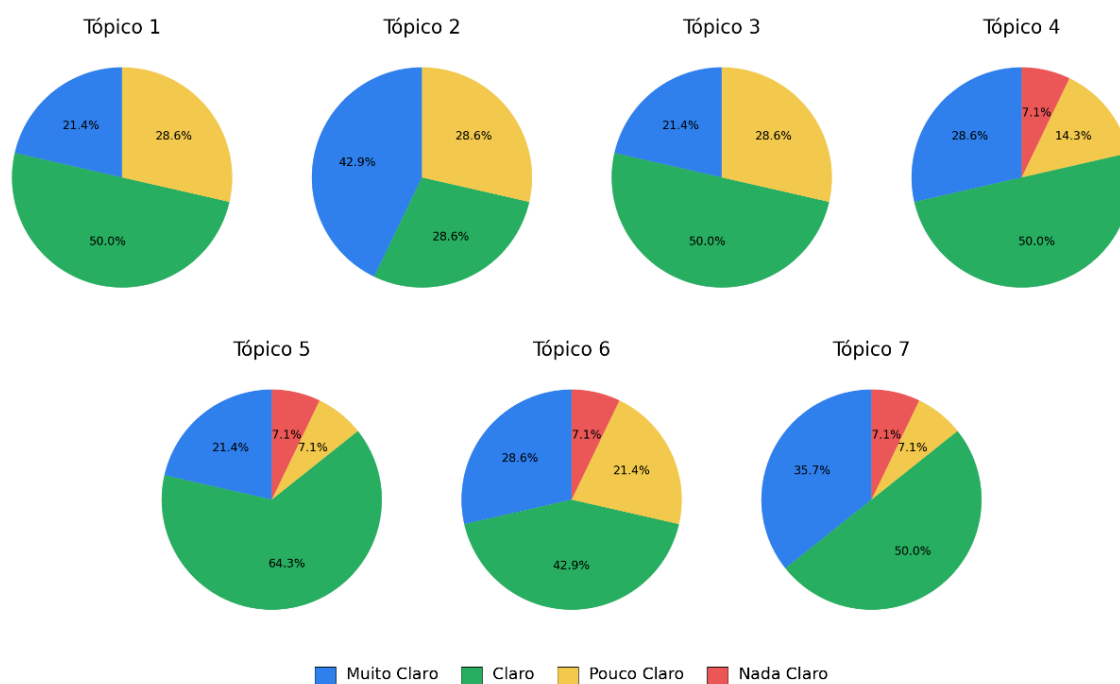


Figura 21: Avaliação da clareza dos tópicos com base nos termos dominantes e incidentes representativos

O tópico 6 apresenta uma boa avaliação do título, indicando uma rotulagem adequada da causa principal, mas uma menor clareza temática, o que pode ser explicado pela natureza técnica e heterogénea dos incidentes representativos, que dificultam uma interpretação imediata. Já o Tópico 2 destaca-se pela forte coerência entre título e conteúdo, refletindo um tema bem delimitado e recorrente, facilmente reconhecido pelos inquiridos. Por sua vez, o Tópico 7 evidencia elevados níveis de clareza, resultado da uniformidade dos incidentes e da simplicidade do mecanismo causal, mesmo não sendo o tópico mais forte ao nível da avaliação do título.

Relativamente à interpretação dos resultados como um todo, a maioria dos inquiridos considerou que não existiam tópicos semelhantes entre si, com 85,7% das respostas afirmativas. Entre os tópicos identificados como semelhantes encontram-se os conjuntos de tópicos 7 e 16, assim como o 10 e 13. No caso do primeiro conjunto, a semelhança deve-se ao facto de ambos os tópicos abordarem ações tardias por parte da equipa, enquanto o segundo conjunto tem em comum o facto de ambos serem relacionados com clientes *undisclosed*. Quanto à existência de tópicos com conteúdo demasiado vago para estarem agrupados, a maioria considerou que não existiam tópicos demasiado vagos para estarem agrupados, com 92,9% das respostas neste sentido. Os tópicos mais identificados como vagos foram o tópico 14 e 15, tendo sido mencionados apenas uma vez por um dos inquiridos. Entre estes 2, destaca-se o tópico 15 por se tratar de um conjunto de incidentes onde os *setups* estavam feitos antes de a equipa iniciar a atividade dos pagamentos naquele sistema, carecendo de uma resposta técnica para o motivo dos incidentes.

Quanto à utilidade de os tópicos melhorarem o entendimento sobre os incidentes dados pelas equipas, a maioria considerou os resultados muito úteis, com metade dos inquiridos, enquanto 35,7% considerou-os úteis e apenas 14,3% a considerarem pouco útil. A maioria dos inquiridos considera que a segmentação por tópicos ajuda a entender as causas dos incidentes como um todo, obtendo-se 71,4% de respostas positivas e 28,6% da população inquirida a responder que ajuda apenas parcialmente.

Da sua experiência, a maioria também considera que os tópicos mais frequentes correspondem às áreas mais críticas de risco, com 78,6% das respostas positivas e 21,4 % das respostas negativas. Na pergunta ‘O modelo revelou alguma causa de risco que, na sua opinião, normalmente passaria despercebida?’ 78,6% dos inquiridos responderam também que não e os restantes responderam que sim, identificando os tópicos 2, 3, 4, 16 e 15 como causas despercebidas, conforme a Figura 22. Com exceção do tópico 15, todos os tópicos foram assinalados por 2 inquiridos, sugerindo um grau de concordância na identificação das causas que, em circunstâncias normais, poderiam passar despercebidas.

É importante destacar que a pergunta seguinte “Se sim, quais?” era de resposta obrigatória sempre que a opção afirmativa fosse selecionada, o que poderá ter contribuído para um maior número de respostas negativas nesta, constituindo uma limitação da precisão das respostas na pergunta ‘O modelo revelou alguma causa de risco que, na sua opinião, normalmente passaria despercebida?’.

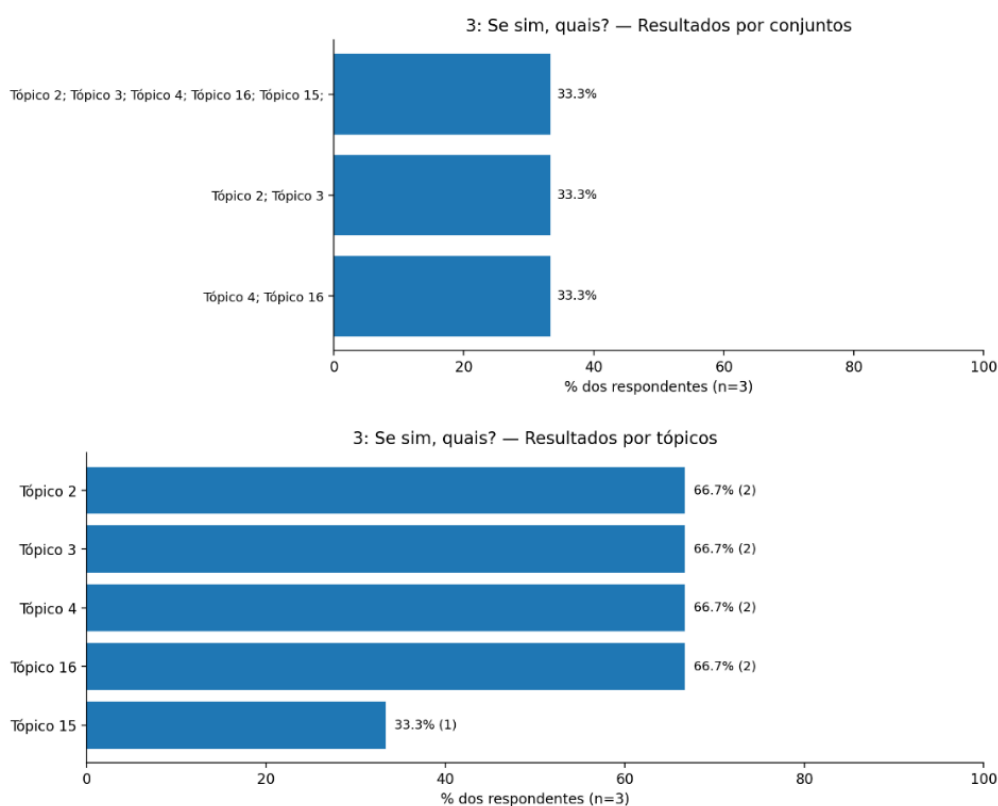


Figura 22: Identificação de causas de risco reveladas pelo modelo “Análise por conjuntos de resposta e por tópicos”

Sobre a utilidade do modelo para a prevenção e melhoria de processos, 92,9% acredita que os tópicos ajudam a discutir ações preventivas quanto aos incidentes e os restantes 7,1% acreditam que o modelo ajuda apenas parcialmente, 92,9% acredita que este tipo de análise ajudaria na identificação precoce de padrões de erro e 7,1% consideram que talvez poderiam ajuda. A maioria da população inquirida vê utilidade na repetição anual deste estudo com dados atualizados, com 92,9% das respostas, com as restantes a considerarem que talvez ajudasse.

Por fim, em termos de perceção geral, os inquiridos avaliam os resultados de forma positiva, com uma média de 4,36 estrelas, com apenas 2 pessoas a atribuir 3 estrelas aos resultados e com a maioria a avaliarem os resultados com 5 estrelas, conforme a Figura 23. Relativamente à questão aberta destinada a comentários adicionais sobre o modelo, apenas dois inquiridos optaram por responder.

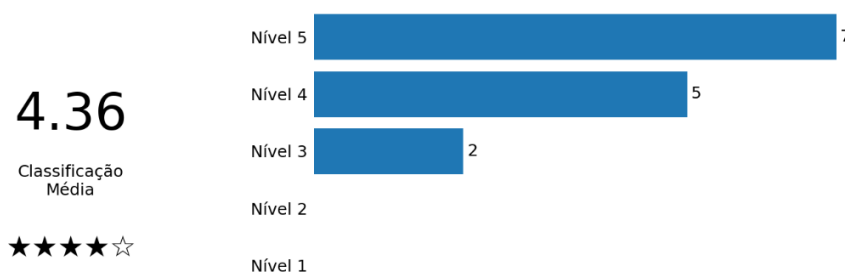


Figura 23: Avaliação global da qualidade da análise - Classificação média atribuída pelos inquiridos

Ambos os comentários convergem na percepção positiva da qualidade global do trabalho, reconhecendo o valor da abordagem utilizada. Ainda assim, é salientada a necessidade de um maior refinamento na interpretação e, sobretudo, na nomeação dos tópicos, de forma a tornar mais explícitas e concretas as causas subjacentes aos incidentes identificados.

5.3 Considerações finais

A validação dos resultados através do inquérito às equipas de *Referential* permitiu avaliar a utilidade prática dos tópicos identificados pelo modelo, através do conhecimento operacional, superando os vieses de interpretações pessoais e da inteligência artificial utilizados neste projeto, complementando estes. Os resultados obtidos demonstram que, apesar da natureza não supervisionada da abordagem, os tópicos gerados são, na sua maioria, compreensíveis, bem rotulados e considerados relevantes para a análise dos incidentes de produção.

A grande maioria dos inquiridos consideraram que a distribuição dos tópicos era útil para a compreensão das principais causas dos incidentes, acreditando também que a segmentação por tópicos ajuda a entender os incidentes e as suas causas de forma geral. Na sua experiência, consideram que os tópicos mais frequentes correspondem às áreas mais críticas. Esta estruturação contribui para reduzir a dependência de análises caso a caso e promove uma visão transversal dos incidentes, alinhada com a percepção das equipas.

Numa perspetiva de negócio, mais de 90% dos inquiridos acreditam que os agrupamentos dos incidentes por tópicos ajudam a discutir ações preventivas e a identificar de forma precoce padrões de erro, com os restantes a acreditarem que talvez possa ser útil, em vez de discordarem. Isto sugere que o modelo pode ser utilizado como um instrumento de apoio à discussão interna e à prevenção de incidentes, permitindo às equipas identificar áreas recorrentes de risco, discutir medidas preventivas de forma informada e acompanhar a evolução dos padrões de erro ao longo do tempo. A elevada aceitação da repetição anual deste tipo de análise reforça a percepção de utilidade do modelo enquanto mecanismo de aprendizagem organizacional contínua.

Apesar de poucas pessoas considerarem que o modelo tenha identificado causas que passam despercebidas, as pessoas que as identificaram acabaram por escolher os mesmos tópicos, reforçando a validade destas respostas. O facto de a resposta afirmativa levar a uma questão de natureza obrigatória que faz o utilizador identificar as causas identificadas nesse sentido, pode ter levado a que alguns utilizadores tenham respondido de forma negativa para evitar responder.

Em síntese, este capítulo procura demonstrar que a aplicação de técnicas de *text mining* a dados não estruturados de incidentes operacionais melhoram significativamente o entendimento das causas dos incidentes de produção em sistemas de pagamentos.

6 Conclusões

Este capítulo final sintetiza os principais resultados obtidos ao longo do trabalho e reflete sobre o seu contributo, limitações e possíveis desenvolvimentos futuros. Para esse efeito, o capítulo estrutura-se em três subsecções: a primeira dedicada aos contributos do estudo, a segunda à identificação das suas principais limitações e a terceira à apresentação de sugestões de melhoria e linhas de investigação futura.

6.1 Contribuições

A presente investigação teve como objetivo central entender como as ferramentas de *text mining* podem ajudar a identificar padrões e melhorar a compreensão das causas dos incidentes operacionais em sistemas de pagamentos, a partir de dados não estruturados extraídos de relatórios de falhas operacionais.

Para isso, foram feitas experiências com abordagens probabilísticas e abordagens com transformadores, fazendo experiências com diferentes números de tópicos e níveis de pré processamento dos dados.

Os resultados obtidos demonstram a eficácia dos modelos de tópicos utilizados em agrupar um conjunto de dados não estruturados, nomeadamente as descrições textuais das causas dos incidentes, em grupos temáticos coerentes e interpretáveis. Em particular destacam-se as abordagens que utilizam o BERTopic, pois conseguiram capturar diferenças entre tópicos através das nuances semânticas.

A qualidade dos resultados e a sua aplicabilidade prática verificou-se através da validação dos resultados feita através de inquérito aos trabalhadores das equipas que fazem parte do ambiente operacional estudado. Os resultados demonstraram que os resultados gerados refletem adequadamente o conteúdo dos incidentes, são compreensíveis e úteis para a discussão interna das causas de risco.

O inquérito demonstrou que os resultados contribuíram para o aumento do conhecimento de forma agregada, onde a maioria acredita que os tópicos mais frequentes correspondem às áreas mais críticas de negócio e, em particular onde existe um consenso sobre os tópicos identificados como causas que poderiam passar despercebidas.

A maioria dos inquiridos considerou estes resultados úteis para a prevenção de risco e identificação precoce das causas de erro e vê potencial na repetição anual do estudo com dados atualizados.

Finalmente, este projeto foi estruturado de forma modular, de modo a garantir que o código possa ser utilizado em análises futuras com alterações mínimas sobre o mesmo, permitindo que o mesmo seja aperfeiçoado e que seja possível aumentar o âmbito da análise tanto a nível temporal como ao nível das equipas analisadas.

6.2 Limitações do projeto

Apesar dos resultados recolhidos, o presente trabalho apresenta um conjunto de limitações e desafios que importa reconhecer em projetos futuros que usem esta metodologia em matéria de incidentes operacionais. Os impactos de alguns dos desafios apresentados foram minimizados de forma considerável através da abordagem exigente e criteriosa deste projeto mas que, mesmo assim, são inerentes à classificação não supervisionada de tópicos pelo que não podem ser completamente superados, merecendo ser incluída neste subcapítulo uma reflexão sobre os mesmos.

Uma das principais limitações associadas à utilização de modelos de tópicos não supervisionados prende-se com a natureza subjetiva da avaliação dos resultados. Com o objetivo de mitigar o impacto da subjetividade dos resultados, o presente estudo desenvolveu uma abordagem que combina a interpretação do autor, a análise feita por uma IA e a validação de resultados feita pela equipa operacional de modo a mitigar o impacto dos vieses associados a cada um destes métodos de avaliação, reforçando a robustez dos resultados obtidos. No entanto, é preciso referir que esta subjetividade não pode ser totalmente eliminada e que deve ser tida em conta ao lidar com estas metodologias.

A fase de pré-processamento revelou-se também um desafio, exigindo um processo de tentativa e erro sobre decisões que dependiam na sua essência do conhecimento de causa. Para além disso, os incidentes foram reportados por diferentes pessoas, desprovidas de um conjunto de regras uniformes para o seu preenchimento, o que resultou em alguns casos em incidentes atribuídos a causas incorretas, criando ruído para a análise. Esta limitação é mais evidente em modelos probabilísticos, exigindo um pré processamento mais exigente. No entanto revelou-se também necessário para os modelos baseados em transformadores, onde alguns dos resultados gerados continham ainda termos que apenas apresentavam ruído para a avaliação dos resultados, mesmo após um nível de pré-processamento que, em teoria, não seria requerido para modelos desta natureza.

Outra limitação relaciona-se com os materiais de visualização disponíveis, que tornam difícil a interpretação dos resultados, mesmo nos modelos que conseguem capturar as nuances semânticas. Dependendo de nuvens de palavras, redes de conceitos e dos termos dominantes de cada tópico aumenta o nível de subjetividade na interpretação e categorização dos tópicos. O BERTopic disponibiliza, como suporte à interpretação, um conjunto de documentos representativos por tópico, o que facilita a compreensão contextual, sobretudo nos tópicos de menor dimensão. Ainda assim, esta informação não permite assegurar, com total certeza, que todos os documentos atribuídos a cada tópico estejam em linha com os temas dominantes identificados.

A principal limitação deste projeto é relativa ao volume reduzido de documentos utilizados nesta metodologia, o que pode afetar a robustez dos resultados e a capacidade de generalização dos tópicos identificados, principalmente nos modelos probabilísticos.

6.3 Melhorias futuras

Com base nos resultados obtidos e nas limitações identificadas, é possível delinear um conjunto de desenvolvimentos futuros que poderão expandir o contributo deste trabalho e a sua utilidade prática.

Para combater a principal limitação deste projeto, propõe-se a expansão do âmbito deste estudo, tanto a nível temporal, analisando um leque temporal maior, como a nível organizacional, expandindo a análise para outras equipas que também trabalhem com pagamentos. Desta forma, é possível melhorar a coerência dos resultados, obtendo tópicos mais consistentes e fidedignos.

Uma recomendação essencial para melhorar a qualidade dos dados passa por estabelecer um conjunto de diretrizes que ajudem a preencher o relatório dos incidentes, de modo a reduzir o impacto do ruído na análise. Para além disso, recomenda-se que sejam criadas duas categorias de comentários, a primeira destinada ao motivo que causou o incidente e a segunda destinada ao processo que levou até à ocorrência do incidente, de modo a evitar que estas duas perspetivas comprometam a qualidade dos resultados por estarem agregadas na mesma variável.

Do ponto de vista metodológico, poderá ser explorado um conjunto de ferramentas complementares que potenciem o valor extraído dos tópicos. Entre estas encontram-se métodos de segmentação que permitiriam retirar informação das outras variáveis disponíveis em conjunto com os tópicos atribuídos a cada documento, como o algoritmo de segmentação categórica, *K-modes*, que permite encontrar padrões entre os tópicos e as restantes variáveis categóricas. Adicionalmente, poderia ser realizada uma análise *Recency, Frequency, Monetary* (RFM), avaliando os tópicos conforme o quão recente foi a última ocorrência, a frequência com que o incidente ocorre e o impacto monetário associado, sendo útil esta análise para reconhecer quais são os incidentes mais sensíveis conforme estas 3 métricas.

Por fim, a repetição anual deste estudo em conjunto com a expansão do âmbito temporal para anos anteriores permitirá uma análise dos incidentes ao longo dos anos, permitindo estudar tendências temporais, a evolução das causas de risco e o impacto das alterações técnicas e processuais ao longo do tempo, permitindo avaliar o impacto das decisões tomadas para prevenir os incidentes com base na análise dos tópicos.

Esta página foi intencionalmente deixada em branco

7 Referências Bibliográficas

- Adams, J. (2002). *Risk: The policy implications of risk compensation and plural rationalities*. UCL Press.
- Aggarwal, C. C. (2023). *Neural networks and deep learning: A textbook* (2nd ed.). Springer. <https://doi.org/10.1007/978-3-031-29642-0>
- Ahadh, A., Binisha, G. V., & Srinivasan, R. (2021). *Text mining of accident reports using semi-supervised keyword extraction and topic modeling*. *Process Safety and Environmental Protection*, 155, 455-465. <https://doi.org/10.1016/j.psep.2021.09.022>
- Al-Ghuribi, S. M., Mohd Noah, S. A., Mohammed, M. A., Tiwary, N., & Saat, N. I. Y. (2024). A comparative study of sentiment-aware collaborative filtering algorithms for Arabic recommendation systems. *IEEE Access*, 12, 174441–174454. <https://doi.org/10.1109/ACCESS.2024.3489658>
- Autoridade Bancária Europeia (EBA). (2017). *Guidelines on incident reporting under PSD2 (EBA-GL-2017-10)*. European Banking Authority. <https://www.eba.europa.eu/guidelines-major-incidents-reporting-under-psd2>
- Banco de Portugal. (2014). Bibliotema: Risco operacional. *Newsletter Biblioteca BdP*, (1), 1–6.
- Bank of England. (2000). *Oversight of payment systems*. Bank of England.
- Basel Committee on Banking Supervision. (2006). *Basel II: International convergence of capital measurement and capital standards: A revised framework—Comprehensive version*. Bank for International Settlements. Retrieved October 13, 2015, from <http://www.bis.org/publ/bcbs128.htm>
- Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent Dirichlet allocation. *Journal of Machine Learning Research*, 3, 993–1022.
- Costa, J. T. M. (2019). *Automatization of incident resolution* [Dissertação de mestrado, ISCTE – Instituto Universitário de Lisboa].
- Committee on Payment and Settlement Systems. (1996). *Settlement risk in foreign exchange transactions*. Bank for International Settlements. <https://www.bis.org/publ/cpss17.pdf>
- Cui, Y., Che, W., Liu, T., Qin, B., Wang, S., & Hu, G. (2020). Revisiting pre-trained models for Chinese natural language processing. arXiv. <https://doi.org/10.48550/arXiv.2004.13922>
- Cui, J., Li, Z., Yan, Y., Chen, B., & Yuan, L. (2023). *Chatlaw: Open-source legal large language model with integrated external knowledge bases* (arXiv:2306.16092). arXiv. <https://arxiv.org/abs/2306.16092>
- Denecke, K., & Paula, H. (2024). Analysis of critical incident reports using natural language processing. *Studies in Health Technology and Informatics*, 313, 1–6. <https://doi.org/10.3233/SHTI240002>
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). *BERT: Pre-training of deep bidirectional transformers for language understanding*. arXiv. <https://doi.org/10.48550/arXiv.1810.04805>

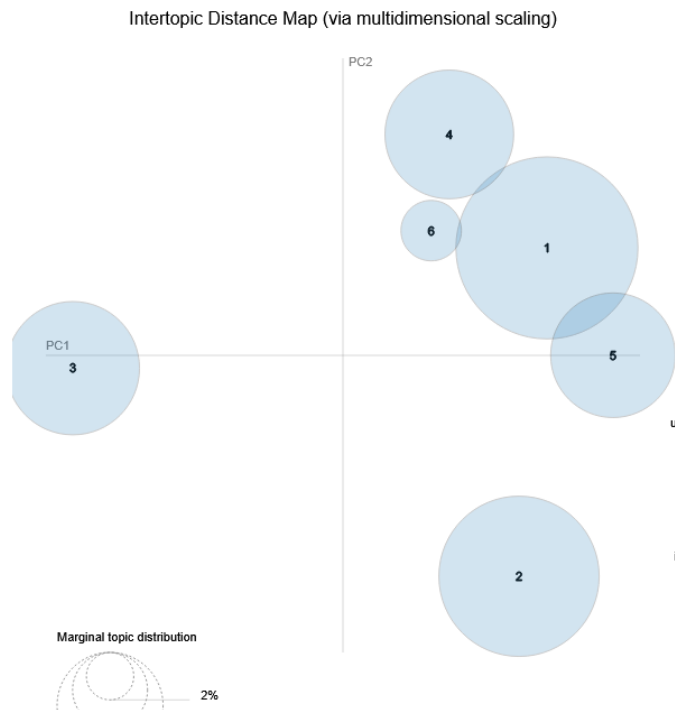
- Dobler, M. C., Garrido, J. M., Grolleman, D. J., Khiaonarong, T., & Nolte, J. (2021). *E-money: Prudential supervision, oversight, and user protection* (Departmental Paper No. DP/2021/027). International Monetary Fund. Retrieved February 10, 2026 from <https://doi.org/10.5089/9781513593401.087>
- Dormosh, N., Abu-Hanna, A., Calixto, I., Schut, M. C., Heymans, M. W., & van der Velde, N. (2024). Topic evolution before fall incidents in new fallers through natural language processing of general practitioners' clinical notes. *Age and Ageing*, 53(2), afae016. <https://doi.org/10.1093/ageing/afae016>
- Fan, W., Wallace, L., Rich, S., & Zhang, Z. (2006). Tapping the power of text mining. *Communications of the ACM*, 49(9), 76–82. <https://doi.org/10.1145/1151030.1151032>
- Ghosh, S., Chen, Y., & Dou, W. (2024). Railroad safety: A systematic analysis of Twitter data. *Case Studies on Transport Policy*, 15, 101154. <https://doi.org/10.1016/j.cstp.2024.101154>
- Giaconia, A., Chiariello, V., & Passarotti, M. (2024). Topic modeling for auditing purposes in the banking sector. In F. Dell'Orletta, A. Lenci, S. Montemagni, & R. Sprugnoli (Eds.), *Proceedings of the Tenth Italian Conference on Computational Linguistics (CLiC-it 2024)* (pp. 1030–1035). CEUR Workshop Proceedings. <https://aclanthology.org/2024.clicit-1.112/>
- Goh, Y. M., & Ubeynarayana, C. U. (2017). Construction accident narrative classification: An evaluation of text mining techniques. *Accident Analysis & Prevention*, 108, 122–130. <https://doi.org/10.1016/j.aap.2017.08.026>
- Grootendorst, M. (2022). *BERTopic: Neural topic modeling with a class-based TF-IDF procedure* (arXiv:2203.05794v1). arXiv. <https://arxiv.org/abs/2203.05794>
- Grootendorst, M. (2024). *BERTopic documentation*. <https://maartengr.github.io/BERTopic/>
- Guagliard, J., Fernandes, N. R., Júnior, A. C., & Júnior, L. R. B. (2016). *Análise de risco parametrizada: Manual prático de gestão de riscos e seguros*. All Print.
- Gutierrez, C. M., & Jeffrey, W. (2006). Minimum security requirements for federal information and information systems (FIPS Publication No. 200). National Institute of Standards and Technology. <https://nvlpubs.nist.gov/nistpubs/fips/nist.fips.200.pdf>
- Heinrich, G. (2006). *Operational risk, payments, payment systems, and implementation of Basel II in Latin America: Recent developments*. Munich Personal RePEc Archive (MPRA) Paper No. 47420. Retrieved from <https://mpra.ub.uni-muenchen.de/47420/>
- Hemrit, W., & Arab, M. B. (2013). The major sources of operational risk and the potential benefits of its management. *Journal of Operational Risk*, 7(4), 71–92. <https://doi.org/10.21314/JOP.2012.115>
- Hotho, A., Nürnberger, A., & Paaß, G. (2005). A brief survey of text mining. *Journal for Language Technology and Computational Linguistics*, 20(1), 19–62. <https://doi.org/10.21248/JLCL.20.2005.68>
- Jorion, P. (2007). *Value at risk: The new benchmark for managing financial risk* (3rd ed.). McGraw-Hill.
- Kang, W. & Cheung, C. F. (2023). Banking System Incidents Analysis Using Knowledge Graph. *International Journal of Knowledge and Systems Science (IJKSS)*, 14(1), 1-23. <https://doi.org/10.4018/IJKSS.325794>

- Khan, M., Khan, S. S., & Alharbi, Y. (2020). Text mining challenges and applications: A comprehensive review. *International Journal of Computer Science and Network Security*, 20(12), 138–148.
- Kim, M., Kim, S., Kim, Y., & Moon, J. (2024). Analyzing the financial impact of ESG news sentiment on ESG finance trends. In *Proceedings of the 2024 International Conference on Platform Technology and Service (PlatCon)* (pp. 95–100). IEEE. <https://doi.org/10.1109/PLATCON63925.2024.10830722>
- Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., Levy, O., Lewis, M., Zettlemoyer, L., & Stoyanov, V. (2019). RoBERTa: A robustly optimized BERT pretraining approach. arXiv. <https://doi.org/10.48550/arXiv.1907.11692>
- Liu, C., & Yang, S. (2022). Using text mining to establish knowledge graph from accident/incident reports in risk assessment. *Expert Systems with Applications*, 207, 117991. <https://doi.org/10.1016/j.eswa.2022.117991>
- Masson, C., & Paroubek, P. (2024). Evaluating topic modeling on imbalanced multi-domain financial corpus for risk extraction. In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)* (pp. 6515–6529). ELRA Language Resource Association.
- McKinsey & Company. (2023). *Response and resilience in operational-risk events*. McKinsey & Company. <https://www.mckinsey.com>
- Miller, P., & Northcott, C. A. (2002). CLS Bank: Managing foreign exchange settlement risk. *Bank of Canada Review*, Autumn, 13–25.
- Mishra, A., Pandey, P., Kumar, A., & Shrivastava, A. (2024). *Temporal analysis of computational economics: A topic modeling approach*. *International Journal of Data Science and Analytics*. <https://doi.org/10.1007/s41060-024-00596-9>
- Ogunleye, B., Maswera, T., Hirsch, L., Gaudoin, J., & Brunsdon, T. (2023). Comparison of topic modelling approaches in the banking context. *Applied Sciences*, 13(2), 797. <https://doi.org/10.3390/app13020797>
- Otal, H. T., Stern, E., & Canbaz, M. A. (2024). *LLM-assisted crisis management: Building advanced LLM platforms for effective emergency response and public collaboration*. 2024 IEEE Conference on Artificial Intelligence (CAI). <https://doi.org/10.1109/CAI59869.2024.00159>
- Operational Risk Exchange (ORX). (2022). *Annual banking operational risk loss data report*. Recuperado em 9 de outubro de 2024, de <https://orx.org/resource/annual-banking-operational-risk-loss-data-report>
- Pedro, A., Pham-Hang, A.-T., Nguyen, P. T., & Pham, H. C. (2022). Data-Driven Construction Safety Information Sharing System Based on Linked Data, Ontologies, and Knowledge Graph Technologies. *International Journal of Environmental Research and Public Health*, 19(2), 794. <https://doi.org/10.3390/ijerph19020794>
- Peters, G. W., Shevchenko, P. V., & Luo, Y. (2018). Rethinking operational risk capital requirements. *Journal of Financial Regulation*, 4(1), 1–34. <https://doi.org/10.1093/jfr/fjx009>
- Royal Society. (1983). *Risk assessment: report of a Royal Society study group*. Royal Society.
- Ruppenthal, J. E. (2013). *Gerenciamento de riscos*. Santa Maria.

- Russo, D., & Stecconi, P. (2011). Operational risks in payment and securities settlement systems: A challenge for operators and regulators. In G. Gregoriou (Ed.), *Operational risk toward Basel III: Best practices and issues in modeling, management, and regulation* (pp. 397-412). Wiley. <https://doi.org/10.1002/9781118267066.ch19>
- Silva, E. S. (2015). *Cadernos Contabilidade e Gestão. Tipologia dos Riscos: Uma Introdução. Vida Económica* pp. 12-17
- Silva, S. A. T. (2018). *Automatization of incident categorization* [Dissertação de mestrado, ISCTE – Instituto Universitário de Lisboa].
- Smetana, M., Salles de Salles, L., Sukharev, I., & Khazanovich, L. (2024). *Highway construction safety analysis using large language models. Applied Sciences*, 14(4), 1352. <https://doi.org/10.3390/app14041352>
- SWIFT. (2024). *Payment operational issues*. SWIFT. Retrieved October 13, 2024, from <https://www.swift.com/pt/node/309355>
- Wang, Z., Li, J., Ma, M., Li, Z., Kang, Y., Zhang, C., Bansal, C., Chintalapati, M., Rajmohan, S., Lin, Q., Zhang, D., Pei, C., & Xie, G. (2024). *Large language models can provide accurate and interpretable incident triage*. 2024 IEEE 35th International Symposium on Software Reliability Engineering (ISSRE). <https://doi.org/10.1109/ISSRE62328.2024.00056>
- Yao, Y., & Li, J. (2022). Operational risk assessment of third-party payment platforms: A case study of China. *Financial Innovation*, 8(1), 19. <https://doi.org/10.1186/s40854-022-00332-x>
- Yang, Z., Wu, Q., Venkatachalam, K., Li, Y., Xu, B., & Trojovský, P. (2022). Topic identification and sentiment trends in Weibo and WeChat content related to intellectual property in China. *Technological Forecasting and Social Change*, 184, 121980. <https://doi.org/10.1016/j.techfore.2022.121980>
- Zhou, M., Kong, Y., & Lin, J. (2022). Financial topic modeling based on the BERT-LDA embedding. 2022 IEEE 20th International Conference on Industrial Informatics (INDIN), 495–500. <https://doi.org/10.1109/INDIN51773.2022.9976145>

8 Anexo

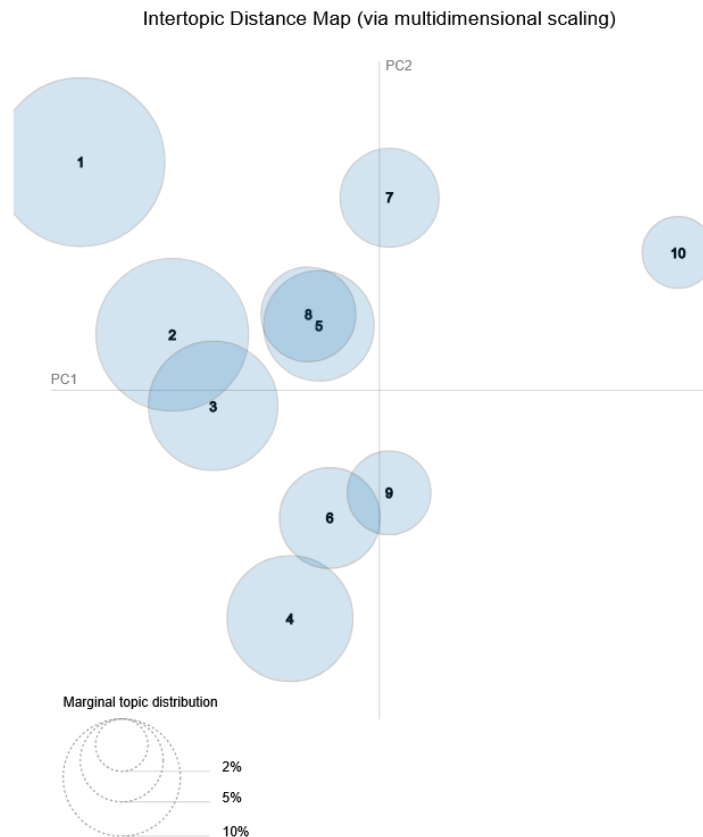
A: Mapa de Componentes Principais e termos mais relevantes - visualização em pyLDAvis, modelo de 6 tópicos



B: Mapa de Componentes Principais e termos mais relevantes - visualização em pyLDAvis, modelo de 6 tópicos.



C: Mapa de Componentes Principais e termos mais relevantes - visualização em pyLDavis, modelo de 10 tópicos.



D: Questionário de Validação Classificação de Causas de Incidentes Operacionais

Questionário de Validação – Classificação de Causas de Incidentes Operacionais

O meu nome é João Saraiva e sou estudante do Mestrado em Business Analytics no Iscte.

Antes de iniciar o preenchimento, leia atentamente a informação abaixo. O tempo estimado de resposta é de **cerca de 10 minutos**.

Este inquérito destina-se exclusivamente às equipas de **CASH SSI** e **FXMM OTC SSI**. O objetivo deste inquérito é **avaliar a qualidade e utilidade dos resultados** obtidos através da análise de tópicos aplicada aos incidentes de 2023. A análise teve como propósito **identificar as principais causas dos incidentes de produção e compreender a sua frequência**, permitindo ainda detetar causas menos evidentes que poderiam passar despercebidas numa avaliação manual.

Para isso, foram analisados os incidentes de ambas as equipas relativas ao ano de 2023. Foram considerados apenas incidentes onde a responsabilidade foi atribuída à equipa de referencial e a causa do incidente foi considerada como Wrong Setup ou Other, desconsiderando pagamentos por atraso ou falta de responsabilidade da equipa. Foram incluídos apenas incidentes de 2023 **atribuídos à equipa de Referencial**, cujas causas estavam classificadas como *Wrong Setup* ou *Other*.

Foram excluídos casos relacionados com **atrasos ou falta de responsabilidade** da equipa, uma vez que o foco está na identificação de causas ligadas ao conteúdo das ações diretas das equipas.

Nota importante:

Algumas causas identificadas podem já não ser aplicáveis atualmente devido a melhorias de sistema, alterações de processo ou evoluções regulatórias. Assim, o objetivo **não é avaliar se a causa ainda existe**, mas sim **validar se os tópicos gerados pelo modelo refletem corretamente padrões reconhecíveis e se a abordagem teria utilidade com dados atualizados**.

Devido ao facto de os incidentes serem **auto reportados**, os resultados dependem da qualidade da informação registada. Podem existir casos com descrições incompletas ou ambíguas, o que limita a precisão da análise e essa perceção também é relevante para este estudo.

A sua participação é voluntária e totalmente anónima. Não serão recolhidos dados que permitam a sua identificação pessoal (como nome ou endereço IP), e todas as respostas serão tratadas de forma confidencial e apenas para fins estatísticos e académicos.

Ao avançar para a próxima página e enviar as suas respostas, declara ter compreendido a natureza deste estudo e aceita participar voluntariamente.

Avaliação Individual dos Tópicos

O modelo identificou um total de 18 tópicos. Para a sua análise e interpretação foram considerados dois elementos fundamentais: os três documentos representativos atribuídos a cada tópico pelo modelo e o conjunto de termos dominantes associados a esse tópico.

Neste inquérito, serão analisados individualmente os 7 tópicos mais relevantes do conjunto de dados, que representam cerca de 50% de todos os incidentes identificados. As causas encontradas dentro dos incidentes são as seguintes:

- Tópico 1: **Notificação de Alert - HKD/CNH**
- Tópico 2: **Alert details incorretos**
- Tópico 3: **SSI_NONE - booking tardio**
- Tópico 4: **Autodebit MT101 – Registo/Aplicação Incorrecta**
- Tópico 5: **Undisclosed Client: Nome apagado do sistema**
- Tópico 6: **Alterações Não Aplicadas (MUR ccy - Bulk FXO)**
- Tópico 7: **Resposta tardia do cliente**
- Tópico 8: **Instruções do Cliente Não Consideradas**
- Tópico 9: **Erros nos Campos 57/58 (Interpretação Incorreta)**
- Tópico 10: **Undisclosed – SSI Errada / BENF Desatualizado (EUR–PLN)**
- Tópico 11: **IBAN Inválido (Dígito em Falta)**
- Tópico 12: **Incidentes Atribuídos Indevidamente**
- Tópico 13: **Undisclosed – Nome Beneficiário Não Atualizado**
- Tópico 14: **Migração: BIC no Campo Errado - Falha de Comunicação**
- Tópico 15: **Old setup: Setup feito antes da equipa tomar responsabilidade**
- Tópico 16: **Broadcast Update: Atualização Não Comunicada / Comunicação Tardia**
- Tópico 17: **Falha na comunicação inter-equipas**

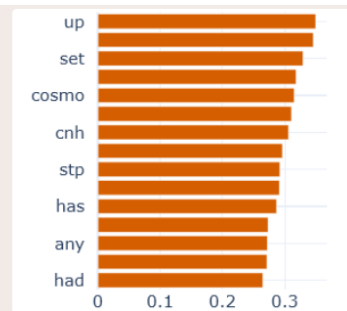
- Tópico 18: **Setups em conformidade**

- Indique o seu User (Ex: E12390).

Nota: As respostas são anónimas. Esta pergunta serve apenas para evitar registos duplicados.

(Introduzir a sua resposta)

Tópico 1: Notificação de Alert - HKD/CNH - 34 Incidentes, 12.4% dos incidentes



Com base na figura apresentada acima, que mostra os termos representativos do tópico e, abaixo, os incidentes representativos associados a esse tópico, responde às seguintes perguntas:

- The REF team follow the process that is to set up before the cut off. That was done as the cut off was at [Hour] . Capacity is the main reason why the team can only manage a % of Urgent cases. However all was done according to the process. If the Cut off is not enough for some teams, maybe it will need to be review so that we can all have a clear view of all the needs.
- We just received the instructions by email on the / / by settlements. However the contact that we had and that we used to get the correct instructions was not correct, ref team did not pay attention that the client name domain was diferente from the actual client. Contact [Client mail] and the email contact [Client mail] Don't belong to the same client. This needs to be shared as we need to have the correct contacts loaded in our system (Local) . Shared Responsibility. Human Distraction.
- We checked with COSMO Support to see if the client alert update was in place and in fact we received the notification for HKD on / / .Please note that HKD had no deals and SSIs setups so we discarded the notification. The alert platform had constraints regarding the non existing CNH notifications however we had a work-around in place for those cases. Whenever we received a HKD (F/X- CASH) notification we also must check manually if this currency received any change in their HKD FX - CASH which corresponds to CNH. And since we had a CNH setup in place from we should have delete or update the SSI depending if it had deals schedule or not.

- **1: Com base na figura apresentada e nos incidentes representativos, considera clara a compreensão do tema abordado pelo tópico?**
 - () Muito claro
 - () Claro
 - () Pouco claro
 - () Nada claro
- **1: O título atribuído aos tópicos reflete corretamente o conteúdo dos incidentes?**
 - () Sim
 - () Não
 - () Talvez

Tópico 2:Alert details incorretos - 31 Incidentes, 11.3% dos incidentes

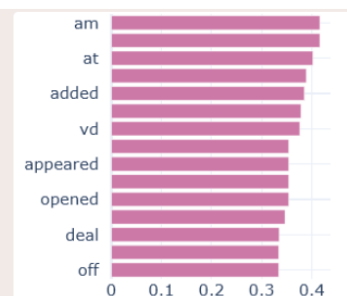


Com base na figura apresentada acima, que mostra os termos representativos do tópico e, abaixo, os incidentes representativos associados a esse tópico, responde às seguintes perguntas:

- Alert details linked and SSI created through Alert on [DATE] with access code [Access Code]. That access code was a perfect match with the short name and nothing indicated it wasn't correct. On [DATE] we received an e-mail saying the alert details were incorrect and that the correct access code was [Access Code]. We amended the code but it wasn't a perfect match with the info on CRDSweb. The short name and code were only updated on CRDSweb later on [DATE]. Everything was done accordingly and on time on our side. Wrong booking
- On the [DAY] of October we received a communication advising that the alert details for this crds were amended by Mumbai team to Access code AB after it was requested by client integration team and we were requested to update the SSIs accordingly. On our end we saw a match so we did the necessary update. As it turns out these access codes were in fact incorrect and on the th of November the client sends another email advising on access code AB which was the one that had been in place since correctly.
- The responsibility for the incident is not for users VBQ and NYSS as they simply corrected the SSI with the most recent information which was provided by the client directly via email. The issue was with the SSI that was previously in place and was created through cosmo by user GSLN and OHTS however they are not the culprits here as the alert details were incorrectly linked on Jun by user CNEW which at the time was part of referential team. Therefore this is our team's responsibility. [DATE] Alert details were updated by user CNEW incorrectly. Wrong alert details linked at [DATE] SSI was created with the instructions that were available on alert at [DATE] SSI was amended after client confirmed the correct instructions via email

- 2: Com base na figura apresentada e nos incidentes representativos, considera clara a compreensão do tema abordado pelo tópico?**
 - ☐ () Muito claro
 - ☐ () Claro
 - ☐ () Pouco claro
 - ☐ () Nada claro
- 2: O título atribuído aos tópicos reflete corretamente o conteúdo dos incidentes?**
 - ☐ () Sim
 - ☐ () Não
 - ☐ () Talvez

Tópico 3:SSI_NONE - booking tardio - 24 Incidentes, 8.7% dos incidentes



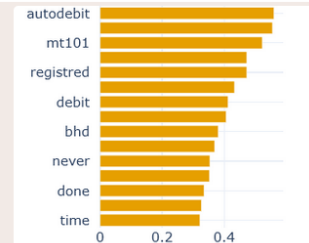
Com base na figura apresentada acima, que mostra os termos representativos do tópico e, abaixo, os incidentes representativos associados a esse tópico, responde às seguintes perguntas:

- The SSI appeared in SSI NONE - Trade Monitoring Queue at [DAY] on the th February. Our Team was not working at this time of day. The AUD currency cutoff was at H AM. There was no option for our team to be able to setup this SSI on time.
- 'SSI_NONE event added on [DAY] at [HOUR]. THB cut-off is at am. Our working hours start at am. There was no way for this deal to be paid on time since we had no visibility before it happened. SSI created during the day. Late booking.'
- "SGD cut-off is at : am. SSI_NONE event was added at : am on the value date. Our working hours begin at am, it would be impossible to create the SSI before the cut-off since we didn't know there was a deal on the day before. Late booking. Early Cutoff."

- 3: Com base na figura apresentada e nos incidentes representativos, considera clara a compreensão do tema abordado pelo tópico?**
 - ☐ () Muito claro
 - ☐ () Claro
 - ☐ () Pouco claro
 - ☐ () Nada claro

- **3: O título atribuído aos tópicos reflete corretamente o conteúdo dos incidentes?**
 - () Sim
 - () Não
 - () Talvez

Tópico 4:Autodebit MT101 - Registo/Aplicação incorreta - 14 incidentes, 5.1% dos incidentes

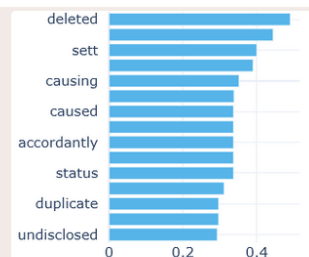


Com base na figura apresentada acima, que mostra os termos representativos do tópico e, abaixo, os incidentes representativos associados a esse tópico, responde às seguintes perguntas:

- SSI done on / before the cut off, but without auto debit since the Iban registration for the auto debit was not done in time by Global Ebics. MT101 Setup
- This is was an issue on our end. Client never requested the BNP SSI to be setup and we never confirmed if autodebit should be applied with them. Moreover we reached out to settlements in order to access if SSIs should be setup as default yes which they denied and said to setup as default no but we still setup both as default yes.
- The SSI that led to payment failure was created on the 1st of December after the client clearly stated that they did not want autodebit MT101. The loro could not be created as default yes because the client also held an account at commerzbank which they only advised was no longer required on the 1st of February already way past the COT. Additionally client only asked us to enable MT101 for the BNP account on the 1st of February.

- **4: Com base na figura apresentada e nos incidentes representativos, considera clara a compreensão do tema abordado pelo tópico?**
 - () Muito claro
 - () Claro
 - () Pouco claro
 - () Nada claro
- **4: O título atribuído aos tópicos reflete corretamente o conteúdo dos incidentes?**
 - () Sim
 - () Não
 - () Talvez

Tópico 5:Undisclosed Client: Nome apagado do sistema - 14 incidentes, 5.1% dos incidentes

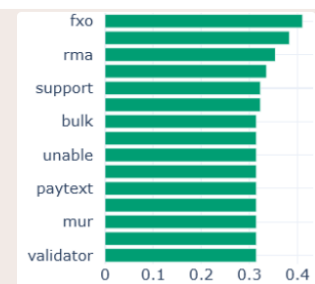


Com base na figura apresentada acima, que mostra os termos representativos do tópico e, abaixo, os incidentes representativos associados a esse tópico, responde às seguintes perguntas:

- SSI_NONE was raised on [DATE] for deal [DEAL N°] SSI team followed up by creating SSI SETT NO: [] , validated on [DATE]. Deal [DEAL N°] paid on [DATE] using SSI SETT NO: [] , causing a cash-out. Following up on non-receipt, the client reached us on VD [DATE] , advising on the wrong payment. . Both the client [CLIENT] and [Front Office Person] @ : followed up by providing the correct SSI. . From here on, SETT NO: [] was created. This setup was correct but was deleted and replaced by SETT NO: [] , an incorrect setup. . Then setup SETT NO:[] was created and validated on[DATE] , which was correctly setup and the deal was correctly paid to the client. Given the above, we consider SETT NO: - as being the root cause of the incident.
- Duplicate as it was reopen The SSI firstly used for payment had an =BENF that was in deleted status which caused the payment rejection. This =BENF was deleted in January and the SSI was not deleted or updated accordantly. The SSI was only amended on the value date April [DAY] at [HOUR] causing the late payment. Undisclosed Setup
- Duplicate as it was reopen The SSI firstly used for payment had an =BENF that was in deleted status which caused the payment rejection. This =BENF was deleted in January and the SSI was not deleted or updated accordantly. The SSI was only amended on the value date April [DAY] at [HOUR] causing the late payment. Undisclosed Setup

- **5: Com base na figura apresentada e nos incidentes representativos, considera clara a compreensão do tema abordado pelo tópico?**
 - () Muito claro
 - () Claro
 - () Pouco claro
 - () Nada claro
- **5: O título atribuído aos tópicos reflete corretamente o conteúdo dos incidentes?**
 - () Sim
 - () Não
 - () Talvez

Tópico 6: Alterações Não Aplicadas (MUR ccy - BULK FXO) - 14 Incidentes, 5.1% dos incidentes



Com base na figura apresentada acima, que mostra os termos representativos do tópico e, abaixo, os incidentes representativos associados a esse tópico, responde às seguintes perguntas:

- On the [DAY] of January all SSIs that were created for MUR currency were amended to include paytext "Payment Related to FX Deal". This was not the fault of creator or validator of the SSI. That being said the communication regarding this change happened in the [DAY] of August. It was initially requested to FXO Admin Support team to amend the SSIs through bulk but given that there was an FXO error soon after they were unable to perform it. As such our team should have actioned this change immediately but we only did it on January th as such we accept this as a team incident.

Nota: Os 3 tópicos representativos continham o mesmo texto.

- **6: Com base na figura apresentada e nos incidentes representativos, considera clara a compreensão do tema abordado pelo tópico?**
 - () Muito claro
 - () Claro
 - () Pouco claro
 - () Nada claro
- **6: O título atribuído aos tópicos reflete corretamente o conteúdo dos incidentes?**
 - () Sim
 - () Não
 - () Talvez

Tópico 7: Resposta tardia do cliente - 13 Incidentes, 4.7% dos incidentes



Com base na figura apresentada acima, que mostra os termos representativos do tópico e, abaixo, os incidentes representativos associados a esse tópico, responde às seguintes perguntas:

- Referential team requested the SSI on [DATE] .We did several chases without reply. The client only replied on [DATE] Client delivered the SSIs after the COT
- Referential team requested the SSI to the client on [DATE], but we didn't include all the contacts we had for this CTPY.Chased done only on [DATE]. This request had several CRDS's and we only received the client reply to CRDS [CRDS CODE] SGH on [DATE].
- Referential SSI requested the SSI to the client on / . The client have alert but the EUR instruction were not available. We did chases without response. The client only replied on [DATE]. Client delivered the SSIs after the COT

- **7: Com base na figura apresentada e nos incidentes representativos, considera clara a compreensão do tema abordado pelo tópico?**
 - () Muito claro
 - () Claro
 - () Pouco claro
 - () Nada claro
- **7: O título atribuído aos tópicos reflete corretamente o conteúdo dos incidentes?**
 - () Sim
 - () Não
 - () Talvez

Avaliação Global dos Resultados

Nesta secção, pretende-se obter a sua perceção global sobre a qualidade e utilidade dos tópicos identificados pelo modelo de análise de incidentes. Com base nesta informação, solicitamos que avalie, de forma geral, a clareza, coerência e relevância dos tópicos gerados, bem como o seu potencial para apoiar a identificação de causas operacionais e a prevenção de novos incidentes.

Interpretabilidade dos Tópicos

- **1: Observou tópicos que parecem demasiado semelhantes entre si?**
 - () Sim
 - () Não
- **1: Se sim, quais? (Escolha múltipla)**
 - Tópico 1
 - Tópico 2
 - Tópico 3
 - Tópico 4
 - Tópico 5
 - Tópico 6
 - Tópico 7
 - Tópico 8
 - Tópico 9
 - Tópico 10
 - Tópico 11
 - Tópico 12
 - Tópico 13
 - Tópico 14
 - Tópico 15
 - Tópico 16
 - Tópico 17
- **2: Identificou algum tópico cujo conteúdo seja demasiado diverso para estar agrupado?**
 - () Sim
 - () Não

- **2: Se sim, quais?** (Escolha múltipla)

- Tópico 1
- Tópico 2
- Tópico 3
- Tópico 4
- Tópico 5
- Tópico 6
- Tópico 7
- Tópico 8
- Tópico 9
- Tópico 10
- Tópico 11
- Tópico 12
- Tópico 13
- Tópico 14
- Tópico 15
- Tópico 16
- Tópico 17

Utilidade dos Tópicos para Compreender os Incidentes

Esta secção procura perceber se os tópicos identificados ajudam, de forma clara e intuitiva, a compreender quais foram as principais causas dos incidentes.

As perguntas seguintes avaliam se a organização dos incidentes por tópicos facilita a análise, a leitura dos dados e a identificação de padrões relevantes, quando comparada com uma análise caso a caso.

- **A distribuição dos incidentes por tópico é útil para a compreensão das principais causas dos incidentes?**
 - () Muito útil
 - () Útil
 - () Pouco útil
 - () Nada útil
- **Considera que a segmentação por tópicos ajuda a entender as causas dos incidentes como um todo?**
 - () Sim
 - () Parcialmente
 - () Não
- **Os tópicos mais frequentes correspondem, na sua experiência, às áreas mais críticas (Desconsiderando atrasos)?**

Nota: Por serem incidentes de 2023, algumas causas poderão estar desatualizadas.

- () Sim
 - () Não
- **O modelo revelou alguma causa de risco que, na sua opinião, normalmente passaria despercebida?**
 - () Sim
 - () Não

- **3: Se sim, quais?** (Escolha múltipla)

- Tópico 1
- Tópico 2
- Tópico 3
- Tópico 4
- Tópico 5
- Tópico 6
- Tópico 7
- Tópico 8
- Tópico 9
- Tópico 10
- Tópico 11
- Tópico 12
- Tópico 13
- Tópico 14
- Tópico 15
- Tópico 16
- Tópico 17

Utilidade para Prevenção e Melhoria de Processos

Nesta secção pretende-se avaliar se os tópicos gerados têm **valor prático e operacional**.

As perguntas procuram perceber se esta forma de análise pode apoiar a discussão de **ações preventivas**, a melhoria de processos existentes e a identificação antecipada de riscos, contribuindo para reduzir a ocorrência de incidentes no futuro.

- **Na sua opinião, os tópicos ajudam a discutir ações preventivas dentro da equipa?**
 - () Sim
 - () Parcialmente
 - () Não
- **Acredita que este tipo de análise ajudaria na identificação precoce de padrões de erro?**
 - () Sim
 - () Talvez
 - () Não
- **Vê utilidade na repetição anual deste tipo de estudos com os dados atualizados?**
 - () Sim
 - () Talvez
 - () Não

Perceção Geral

- **Como avalia globalmente a qualidade da análise produzida?**

Escolher 0 a 5 estrelas

- **Têm comentários adicionais sobre os resultados ou sugestões para melhorar o modelo?**

Introduza a sua resposta

Link: [Questionário de Validação – Classificação de Causas de Incidentes Operacionais – Preencher o formulário](#)