

# Guia prático para implantar chatbot

com 100% de sucesso



Olá! Eu sou o  
**Milbot**. Como  
posso te ajudar?



  
MILVUS

# { SUMÁRIO

|    |                                                                           |    |
|----|---------------------------------------------------------------------------|----|
| 01 | ANTES DO BOT: POR QUE A MAIORIA FALHA<br>(E COMO NÃO FALHAR)              | 03 |
| 02 | ARRUME A CASA: DADOS + BASE DE CONHECIMENTO<br>(TOKENIZAÇÃO SEM MISTÉRIO) | 05 |
| 03 | IA NOS BASTIDORES & POLÍTICA DE RESPOSTA SEGURA<br>(ANTI-ALUCINAÇÃO)      | 10 |
| 04 | PLANO DE IMPLANTAÇÃO EM 4–6 SEMANAS                                       | 15 |
| 05 | MÉTRICAS QUE IMPORTAM                                                     | 18 |
| 06 | OPERAR, EVOLUIR E INTEGRAR<br>(DO TICKET AO CONHECIMENTO + CHATBOT 2.0)   | 22 |
| 07 | COMO O MILVUS APLICA IA NA PRÁTICA<br>PARA TRANSFORMAR O SERVICE DESK     | 27 |
|    | <b>Anexos</b>                                                             |    |
|    | A.CHECKLIST DE SUCESSO (1 PÁGINA)                                         | 30 |

# 01

## ANTES DO BOT: POR QUE A MAIORIA FALHA (E COMO NÃO FALHAR)

A maioria dos projetos de chatbot não fracassa por “falta de IA”, e sim por falta de base: dados espalhados entre canais, respostas contraditórias, artigos desatualizados e ausência de regras claras.

Sem uma governança mínima (quem publica, quem revisa, o que entra/sai da base), qualquer modelo por melhor que seja acaba errando o alvo. Este capítulo coloca todo mundo na mesma página antes de falar de bot.

### Por que projetos escorregam?

1. **Dados desorganizados:** tickets, e-mails e WhatsApp sem unificação, sem deduplicação e sem padrão de termos.
2. **Base frágil:** artigos sem dono, versões conflitantes, partes críticas sem homologação.
3. **Políticas ausentes:** o bot responde “no improviso”, sem regra de “não sei” nem fluxo de fallback (plano B seguro quando a resposta não é confiável).
4. **Métricas vagas:** ninguém define linha de base para % resolvidos pelo bot, MTTR, FCR e CSAT. Fica impossível provar melhoria.

### Glossário rápido:

- **FCR — First Contact Resolution (Resolução no Primeiro Contato):**  
% de casos resolvidos no primeiro atendimento.
- **MTTR — Mean Time To Resolution (Tempo Médio para Resolução):**  
tempo médio do “abrir” ao “resolver”.
- **CSAT — Customer Satisfaction Score (Índice de Satisfação do Cliente):**  
nota de satisfação pós-atendimento.
- **Fallback:**  
plano B quando o bot não tem confiança para responder; pode abrir ticket, transferir para humano ou pedir mais dados.



# ANTES DO BOT:

## POR QUE A MAIORIA FALHA (E COMO NÃO FALHAR)

01

### O que significa “100% de sucesso”?

“100% de sucesso” não é o bot acertar tudo; é operar com previsibilidade e segurança.

Significa responder com base homologada quando há cobertura; assumir limites (dizer “não sei”) quando não há; acionar fallback de forma elegante; proteger dados (LGPD) com rastreabilidade; e medir para melhorar continuamente.

Em outras palavras: qualidade consistente hoje e melhor do que ontem, sem sustos.

### Para saber se você está no caminho certo, existem alguns sinais práticos, como por exemplo:



Você sabe o que entra na base e quem aprova (e consegue apontar os donos de cada tema).



Perguntas repetidas já estão mapeadas e têm respostas padronizadas e fáceis de achar.



Quando o bot não sabe, ele não inventa: segue um fallback claro, e esse fluxo é medido.

**Fale com um de nossos Especialistas**





# ARRUME A CASA: DADOS + BASE DE CONHECIMENTO (TOKENIZAÇÃO SEM MISTÉRIO)

02

## Como funciona a fundação de dados?

Antes do bot falar com alguém, organize a informação. Isso significa mapear e unificar tudo o que hoje responde dúvidas do cliente (**tickets, e-mails, WhatsApp, CRM, FAQs, manuais**) em um formato único, sem duplicidades e com linguagem padronizada.

Quando o conteúdo fica limpo e consistente, a IA encontra o que precisa e você consegue medir e auditar o que foi usado para responder.

## O que fazer:

1. **Mapeie fontes e campos mínimos:** Liste de onde virão as respostas (ex.: "Tickets", "E-mail", "WhatsApp", "CRM"). Defina um modelo único com campos como: Título, Problema, Solução passo a passo, Pré-requisitos, Exceções, Links úteis, Data/versão, Dono do conteúdo.
2. **Deduplicação e padronização de termos:** Uma conteúdos repetidos, escolha uma versão "canônica" e crie um glossário vivo de termos oficiais e sinônimos (ex.: "2ª via de boleto" = "segunda via"). Isso reduz ruído e ambiguidade para pessoas e IA.
3. **Mascaramento de PII e LGPD desde o início:** PII – Personally Identifiable Information (dados pessoais identificáveis) não deve ir para a base bruta: mascare nomes, CPFs, e-mails, telefones, endereços (ex.: João Silva → [NOME], 123.456.789-09 → [CPF]). Registre quem alterou o quê e quando (trilha de auditoria) e nomeie responsáveis por cada tema.
4. **Critérios de publicação (governança mínima):** Deixe claro o que entra (fontes oficiais, nível de detalhe) e quem aprova. Artigos precisam de dono e versão. Sem isso, a base "apodrece" rápido.

## Resultado esperado:

conteúdo limpo, rastreável e pronto para consulta da IA,  
com responsáveis definidos e histórico de mudanças.

# ARRUME A CASA: DADOS + BASE DE CONHECIMENTO (TOKENIZAÇÃO SEM MISTÉRIO)

## Tokenização (chunking) sem mistério

Em vez de guardar textos longos, divida cada artigo em trechos curtos (chunks) e etiquete cada trecho. Assim, a IA recupera só a parte relevante, o que melhora a precisão e reduz o custo de processamento.

### → Como quebrar:

uma ideia por trecho (ex.: 120–300 palavras), com etiquetas como Assunto, Produto/Área, Tipo (procedimento, FAQ, política), Versão.

### → Por que funciona:

a pergunta do usuário é comparada com trechos (não com o artigo inteiro), aumentando a chance de bater exatamente no conteúdo certo.

### → Rastreabilidade:

toda resposta da IA deve registrar quais trechos/artigos foram usados (IDs/versões). Se o cliente marcar “não resolveu”, você sabe exatamente o que revisar.

### Exemplo de metadados por trecho (mínimo viável):

```
{ artigo_id, versao, trecho_id, assunto, palavras-chave, data_publicacao, dono }
```

## Checklist “pronto para indexar”

- ✓ Fontes unificadas no modelo único.
- ✓ Duplicatas removidas e termos padronizados (glossário).
- ✓ PII mascarada e LGPD ok (com auditoria).
- ✓ Artigos “chunkados” e etiquetados.
- ✓ Donos e versões definidos.
- ✓ Log configurado para citar trechos usados em cada resposta.

Com a fundação de dados pronta e tokenização aplicada, você tem uma base confiável e auditável. No próximo trecho do capítulo, vamos transformar essa base em conteúdo pronto para IA (estrutura de artigos, critérios de homologação e manutenção).



# ARRUME A CASA: DADOS + BASE DE CONHECIMENTO (TOKENIZAÇÃO SEM MISTÉRIO)

# 02

## BASE DE CONHECIMENTO PRONTA PARA IA

A base de conhecimento é o manual oficial do seu atendimento. Quando ela está bem estruturada e governada, a IA consegue recuperar o trecho certo, dar respostas consistentes e aprender com cada interação.

## Estruturade artigos (template recomendado)

### 1. Cada artigo deve seguir um padrão simples, previsível e fácil de auditar.

- Cabeçalho (metadados)
- Título atômico (1 problema = 1 solução)
- Assunto / Produto / Área
- Palavras-chave e sinônimos (ex.: "2ª via", "segunda via", "boleto")
- Nível (N1/N2), Ambiente (web/app), Plataforma (Windows/Linux/etc.)
- Dono (responsável pelo conteúdo), Versão, Data de revisão/expiração
- Estado: Rascunho → Homologado → Obsoleto (com redirecionamento)

### Dica:

se o artigo for longo, quebre em trechos (chunks) por subtópico. Cada chunk tem um título claro e metadados próprios (assunto/tags).



# ARRUME A CASA: DADOS + BASE DE CONHECIMENTO (TOKENIZAÇÃO SEM MISTÉRIO)

02

2.

## Padrões de escrita (para pessoas e IA)

- **Clareza > Jargão:** frases curtas, termos padronizados (use o glossário).
- **1 ideia por parágrafo;** 1 ação por passo.
- **Chame o objeto pelo nome oficial** (ex.: “2ª via de boleto”, não “segundo boleto”).
- **Títulos e subtítulos específicos:** “Reset de senha no app Android” (melhor que “Senha”).
- **Sinônimos/campos** “Também conhecido como” para melhorar a busca.
- **Nada de PII** (dados pessoais) em exemplos/prints.

3.

## Homologação e versionamento

- **4 olhos:** autor (área dona) + revisor técnico/operacional.
- **Checklist de homologação:** precisão validada, cobertura de exceções, LGPD/PII verificados, “Definition of Done” definida.
- **Versionamento visível:** versão, data, quem alterou, motivo da alteração, link para o ticket que originou a mudança.
- **Descontinuação:** marcar como obsoleto e redirecionar para artigo substituto (evita links quebrados).

4.

## O que entra / o que não entra

### Entra

- Procedimentos repetitivos e de alto volume.
- Políticas oficiais e respostas homologadas por área responsável.
- Mensagens padrão de abertura/encerramento e templates de comunicação.

### Não entra

- Opiniões, “dicas de corredor”, hipóteses não testadas.
- Prints com PII (troque por prints limpos ou passo a passo textual).
- Segredos de acesso (credenciais) ou dados sensíveis.





# ARRUME A CASA: DADOS + BASE DE CONHECIMENTO (TOKENIZAÇÃO SEM MISTÉRIO)

02

5.

## Como otimizar busca e recuperação pela IA

- **Títulos atômicos e específicos** (problema + contexto).
- **Assuntos/Tags coerentes**; mantenha um glossário de sinônimos.
- **Chunking**: trechos curtos com um mini-título descritivo (“Gerar 2ª via pelo portal web”).
- **Primeira frase** do chunk deve “dizer ao que veio” (resumo em uma linha).
- **Metadados mínimos por chunk**: {artigo\_id, versao, trecho\_id, assunto, palavras-chave, dono, data}.
- **Rastreabilidade**: cada resposta do bot aponta quais trechos sustentaram a resposta.

### Exemplo de título

- **Ruim**: “Fatura”
- **Bom**: “Gerar 2ª via de boleto no portal web (Pessoa Jurídica)”

6.

## Auditoria e manutenção contínua

- **Trilha de auditoria sempre ativa**: quem alterou, quando e por quê.
- **Revisão periódica**: sugerido trimestral para top 20 artigos e semestral para o restante.
- **Expiração programada**: defina data de revisão; artigos vencidos entram em fila de atualização.
- **Donos por tópico**: cada assunto tem um responsável claro (não “área inteira”).
- **Ciclo ticket → conhecimento**: todo ticket recorrente/complexo vira novo artigo (com link para o ticket de origem).

Com essa base homologada, versionada e chunkada, a IA recupera com precisão, o time confia no conteúdo e você tem controles para manter tudo atualizado sem virar um gargalo.

# IA NOS BASTIDORES & POLÍTICA DE RESPOSTA SEGURA (ANTI-ALUCINAÇÃO)

03

## IA DE BASTIDORES (ANTES DE IR A PÚBLICO)

Antes de colocar um bot para falar com o cliente, **use a IA internamente para acelerar o time e validar o processo**. Aqui, a IA atua como copiloto do analista: lê o histórico, resume o chamado, sugere um rascunho de abertura/encerramento, indica possíveis respostas utilizando a base de conhecimento homologada e ajuda na classificação.

Nada é enviado ao cliente automaticamente; o analista revisa, edita (se quiser) e decide. **Isso reduz atrito, dá previsibilidade e evita o cenário em que o modelo “inventa” sem garantia.**

A qualidade das sugestões não deve ser “no olho”. Adote um limiar de assertividade: uma checagem entre **pergunta ↔ resposta ↔ trechos da base**. Quando a pontuação fica acima do limiar (ex.: 70%), o rascunho é apresentado como “apto”; quando fica abaixo, o sistema sinaliza “precisa de revisão” e não recomenda envio automático.

Essa verificação obriga a resposta a apontar quais trechos (chunks) a sustentam, preservando a rastreabilidade e facilitando auditoria e melhoria contínua.

Perguntas ambíguas são outra fonte clássica de erro. Em vez de forçar uma resposta fraca, a **IA pode propor reformulações com base na base de conhecimento e nos termos mais frequentes** (ex.: “Você quer 2ª via do boleto ou alterar o vencimento?”).

No piloto de bastidores, essas reformulações servem para treinar o time e mapear intenções; quando o canal for aberto ao público, elas viram opções claras para o cliente escolher, aumentando a precisão e reduzindo re-trabalho.

Operacionalmente, vale padronizar tom de voz (objetivo, cordial, sem jargão), limites de resposta (máx. de caracteres/palavras) e modelos de abertura/encerramento. Esse “guarda-corpo” deixa as sugestões consistentes e fáceis de aprovar.

Em paralelo, rode um 80/20: **identifique as 10–20 perguntas que concentram a maior parte do volume e garanta que existem artigos bons (e chunkados) para cada uma**. É nessas frentes que os ganhos aparecem rápido (primeira resposta mais veloz, registro mais limpo e menos idas e vindas) enquanto você prova valor, treina o time e coleta evidências para o business case.

# **IA NOS BASTIDORES & POLÍTICA DE RESPOSTA SEGURA (ANTI-ALUCINAÇÃO)**

# 03

## **Saídas mensuráveis em 2 semanas** (sugestão de leitura executiva):

- Tempo até a 1ª resposta ↓ com rascunhos de abertura e resumos automáticos.
- Qualidade do registro ↑ (campos completos e consistentes, menos retrabalho).
- % de sugestões “aptas” (acima do limiar) e principais lacunas da base mapeadas.
- Top intenções 80/20 definidas e priorizadas para a próxima fase.

## **Boas práticas de configuração** (mínimo necessário):

- Ativar resumo + rascunhos (abertura/encerramento) + sugestões de resposta sempre ancoradas em trechos da base de conhecimento.
- Definir limiar de assertividade e exigir referência aos chunks nas sugestões.
- Habilitar reformulações para casos ambíguos e registrar qual opção foi escolhida.
- Logar tudo: quem aprovou, o que foi usado na resposta e por que (para aprender rápido e com segurança).

Com a IA nos bastidores, você ganha velocidade sem perder controle, transforma cada atendimento em dado útil e chega ao “abrir ao público” com base, processos e métricas testados. Na próxima parte deste capítulo, formalizamos a política de resposta segura (anti-alucinação) e os fluxos de fallback.

**Agende uma Demonstração Gratuita do Milvus**

# IA NOS BASTIDORES & POLÍTICA DE RESPOSTA SEGURA (ANTI-ALUCINAÇÃO)

# 03

## POLÍTICA DE RESPOSTA SEGURA (ANTI-ALUCINAÇÃO) E FALLBACK (PLANO B SEGURO)

### Regra de ouro:

o bot só responde com conteúdo homologado na base de conhecimento. Se faltar cobertura ou confiança, assume “não sei” (isso é qualidade, não erro) e aciona fallback.

### Quando NÃO responder (gate de segurança)

- ✗ Assertividade abaixo do limiar (ex.: < 70%) entre pergunta ↔ resposta ↔ trechos da base.
- ✗ Pergunta ambígua (múltiplas intenções possíveis) ou falta dado crítico.
- ✗ Tema sensível / fora de escopo (dados pessoais/financeiros, credenciais, políticas internas).
- ✗ Artigo desatualizado ou sem homologação.

Toda resposta aprovada deve citar quais trechos (chunks/IDs) a sustentam. Sem referência → não publica.

# IA NOS BASTIDORES & POLÍTICA DE RESPOSTA SEGURA (ANTI-ALUCINAÇÃO)

# 03

## Fluxos de fallback (plano B seguro):

**Abrir ticket, Transferir para humano  
ou Coletar dados / desambiguação.**

### 1.

#### **Abrir ticket**

- **Coleta campos mínimos** (assunto, descrição, contato, anexos) → gera protocolo → envia para a fila correta.
- **Mensagem-padrão (curta):** “Não tenho informação suficiente para responder com segurança agora. Abri o chamado nº {protocolo} para o nosso time continuar por aqui.”

### 2.

#### **Transferir para humano**

- **Encaminha com contexto:** pergunta do cliente, rascunho (se houver), trechos consultados e motivo da transferência.
- **Mensagem-padrão:** “Vou te transferir para um atendente que continua a partir do ponto em que paramos.”

### 3.

#### **Coletar dados / desambiguação**

- **Faz pergunta objetiva** ou oferece reformulações (2–3 opções) com base na base de conhecimento.
- **Mensagem-padrão:** “Para te ajudar melhor, escolha uma opção: [2ª via de boleto] [Alterar vencimento] [Outra dúvida].”

**Padrão recomendado:** tentar desambiguação → se mantiver baixa confiança, abrir ticket; transferência para casos urgentes/sensíveis.

#### **Validação de perguntas (classificar a intenção)**

- Classifique a entrada como Pergunta, Reclamação, Solicitação administrativa (ex.: boleto) ou Incidente técnico.
- Use a classe para rotear corretamente (fila/regra) e para oferecer opções de desambiguação coerentes.
- Registre a intenção escolhida pelo cliente para treinar base e métricas.



# IA NOS BASTIDORES & POLÍTICA DE RESPOSTA SEGURA (ANTI-ALUCINAÇÃO)

03

## Mensagens-padrão (curtas) para o bot

- “Não sei” com elegância: “Prefiro não arriscar uma resposta incorreta. Posso abrir um chamado ou você prefere falar com um atendente?”
- Desambiguação: “Você se refere a 2ª via de boleto ou alterar vencimento?”
- Confirmação de ticket: “Pronto! Chamado {protocolo} aberto. Avisaremos aqui quando houver atualização.”

## Benefícios (por que adotar essa política)

- ✓ Menos risco (zero “alucinação crítica”), mais confiança do usuário e da operação.
- ✓ Auditoria facilitada: cada resposta tem rastro (trechos/IDs, quem aprovou).
- ✓ Melhoria contínua: motivos de fallback viram pauta de novos artigos e ajustes de processo.

## Métricas de saúde (acompanhe de perto)

- ➔ Taxa de fallback (deve cair com a evolução da base).
- ➔ % de assertividade acima do limiar (tendência estável/ascendente).
- ➔ Zero alucinação crítica + tempo até handoff (quando transfere).
- ➔ CSAT pós-fallback e reabertura de casos (devem cair).

# PLANO DE IMPLANTAÇÃO EM 4-6 SEMANAS

04

A ideia é colocar o chatbot no ar com segurança e previsibilidade, começando pequeno, fora do horário comercial, e avançando por etapas claras. Em cada passo, medimos resultados e só seguimos adiante quando estiver funcionando bem.

## Semanas 1-2

### Piloto:

#### Abertura de tickets (off-hours)

**Comece simples:** o bot só abre chamados. Ele coleta os campos mínimos (assunto, descrição, contato/anexos), gera o protocolo e direciona para a fila certa.

Aqui, a Base de Conhecimento já ajuda a qualificar a coleta de dados (perguntas objetivas) e, se faltar informação, o bot assume “não sei” e segue para fallback (abrir ticket, transferir para humano ou pedir mais dados).

**O que observar:** chamados entram completos, tempo até a primeira resposta humana e por que o bot acionou fallback. Com isso, você mapeia as 10-20 intenções mais frequentes (regra do 80/20) e as lacunas da Base de Conhecimento.

## Semana 3

### Classificação assistida

#### (off-hours)

**Adicione a classificação:** o bot sugere categoria e prioridade com base na pergunta e em trechos da Base de Conhecimento; um humano confirma em segundo plano. Você mede acerto das sugestões e onde a Base ainda está fraca. Ajuste termos, sinônimos e títulos para a recuperação de conteúdo ficar mais certa.

# **PLANO DE IMPLANTAÇÃO** **EM 4–6 SEMANAS**

## **Semana 4** **Respostas N1** **(off-hours, parte do tráfego)**

Agora inclua um conjunto pequeno de dúvidas simples e repetidas (N1) às 10–15 mais comuns. O bot só responde quando encontra trechos homologados da Base e a assertividade (a comparação entre pergunta, resposta e trechos usados) fica acima do limiar combinado (ex.: 70%).

Se não atingir, ele não arrisca: usa “não sei” e aciona o fallback. O que observar: % resolvido pelo bot, tendência de assertividade, CSAT pós-atendimento e principais motivos de fallback (cada motivo vira pauta de melhoria na Base).

## **Semanas 5–6** **Horário comercial** **e expansão gradual**

Se os números se mantiverem bons por pelo menos duas semanas, passe do off-hours para o horário comercial em 30–50% do tráfego. Estabilizando, avance para 100%. Depois, repita o mesmo processo em outros canais (site → WhatsApp/Instagram).

### **Regra simples para Go/No-Go:**

avance quando a assertividade estiver no limiar ou acima, a cobertura da Base para as intenções do piloto estiver em torno de 80%, a taxa de fallback for controlada e não houver alucinação crítica nem incidentes de LGPD/PII. Se não bater, segure a etapa atual, melhore Base/política e meça de novo.

# PLANO DE IMPLANTAÇÃO EM 4–6 SEMANAS

04

## Se algo sair do esperado

Tenha um plano de reversão simples: desative temporariamente as respostas NI, mantenha Abertura e Classificação, redirecione tudo para humano com contexto, congele novas intenções e faça uma revisão rápida (24–48h) da Base, prompts e critérios. Depois, retome a evolução.

## Dimensione o esforço (S/M/L) desde o início

Para evitar surpresas, classifique o tamanho do projeto com quatro entradas fáceis: tickets por mês, número de intenções iniciais, cobertura atual da Base de Conhecimento e número de integrações. Revise essa fotografia a cada checkpoint de Go/No-Go: ela orienta prazos, pessoas e prioridades.

## Como comunicar com o time (e ganhar adesão)

- 1 | Transparência traz colaboração.** Explique que a IA é um copiloto para tirar o peso do repetitivo e que o bot só responde com conteúdo homologado da Base.
- 2 | Mostre as metas públicas** (assertividade, cobertura, CSAT) e convide o time a marcar “resolveu/não resolveu” e a apontar lacunas. Isso vira novo conteúdo!
- 3 | Reserve uma reunião curta semanal** (15 minutos) para revisar casos reais e celebrar pequenas vitórias (queda de MTTR, menos reaberturas).

Em toda comunicação, mantenha linguagem simples, tom cordial, limites de tamanho de resposta e atenção a LGPD/PII (nada de dados pessoais em prints ou exemplos).

## Medir é o que separa “promessa” de resultado.

**Antes de abrir o bot ao público,** defina uma linha de base (como estamos hoje) e metas para o período do piloto. Aqui estão as métricas que realmente importam, como calcular e como ler em linguagem simples.

### NÚCLEO DO PILOTO (O QUE NÃO PODE FALTAR)

% resolvidos pelo bot — parcela dos atendimentos concluídos sem intervenção humana.

- ➔ **Como calcular:** tickets elegíveis resolvidos pelo bot ÷ total de tickets elegíveis do piloto.
- ➔ **Leitura:** sobe quando a Base de Conhecimento cobre bem as intenções; cai quando falta conteúdo ou há muito caso fora de escopo.

### MTTR — Mean Time To Resolution (Tempo Médio para Resolução)

- ➔ **Como calcular:** média de (hora de resolução - hora de abertura).
- ➔ **Leitura:** meça geral e do bot; a tendência deve cair à medida que as dúvidas simples migram para o N1 automatizado.

### FCR — First Contact Resolution (Resolução no Primeiro Contato)

- ➔ **Como calcular:** casos resolvidos no primeiro contato (sem transferir/reabrir) ÷ total de casos.
- ➔ **Leitura:** indica eficiência e clareza; melhora quando o bot oferece respostas completas ou coleta dados certos de primeira.



## CSAT — Customer Satisfaction Score (Satisfação do Cliente)

- **Como calcular:** média das notas pós-atendimento (ex.: 1–5).
- **Leitura:** monitore por canal e por tipo de atendimento (bot vs. humano). Quedas súbitas pedem revisão de mensagens e conteúdo.

## CONTROLE DE RISCO E QUALIDADE

**Assertividade (%) —** quanto bem a resposta bate com a pergunta e os trechos da Base de Conhecimento.

- **Como usar:** defina um limiar mínimo (ex.: 70%). Abaixo disso não responde; aciona “não sei” + fallback.
- **Leitura:** se a média cai, falta artigo bom, o título está genérico ou a pergunta chega ambígua.

## Taxa de escalonamento

quanto o bot precisa passar para humano.

- **Como calcular:** casos transferidos/abertos ÷ casos que o bot recebeu.
- **Leitura:** deve cair com a evolução da Base; se sobe, liste motivos e transforme em novos artigos.

## Tempo de registro

quanto o agente leva para registrar/estruturar um chamado.

- **Como calcular:** do início do registro até salvar com campos completos.
- **Leitura:** cai quando usamos resumos automáticos e rascunhos; é um bom indicador de produtividade do time.

## Qualidade dos dados do ticket

completude e consistência dos campos.



### Como medir (índice simples):

% de tickets com campos obrigatórios completos, sem ambiguidades, com categoria correta e link para artigo/trecho usado.



### Leitura:

dados ruins viram métricas ruins. Corrija o formulário e padronize termos.

## COMO INSTRUMENTAR

(SEM COMPLICAR)

- **Rastreabilidade:**  
toda resposta do bot deve citar os trechos (IDs/versões) da Base usados.
- **Amostragem quinzenal:**  
revise 15–20 casos reais (QA) para checar clareza, tom e aderência à Base.
- **Linha de base:**  
colete 2–4 semanas de dados antes do piloto para comparar.

[Clique aqui e veja o Milvus em ação](#)

## DASHBOARDS SIMPLES (LEITURA EXECUTIVA)

1. Saúde diária: % resolvidos pelo bot, assertividade média, taxa de escalonamento, incidentes (alucinação crítica = 0).
2. Tendência quinzenal: MTTR, FCR, CSAT e tempo de registro (todas com setas de tendência).
3. Diagnóstico: top motivos de fallback e lacunas da Base (viram pauta de novos artigos).

## CADÊNCIA E DECISÕES

1. Revisão quinzenal: olhar o painel, decidir Go/No-Go da próxima etapa e priorizar 3 ações (ex.: criar 5 artigos, ajustar 3 títulos, refinar 2 mensagens do bot).
2. Alertas práticos: assertividade < limiar por horas seguidas, CSAT em queda relevante, pico de reaberturas segura a expansão, melhore a Base e reteste.
3. Transparência: compartilhe as métricas com o time; peça "resolveu/não resolveu" em cada atendimento para retroalimentar o ciclo.

### Regra de ouro:

se você consegue medir bem, consegue melhorar rápido. Essas métricas mostram onde atuar, comprovam valor do piloto e dão segurança para expandir canais e horários.



# OPERAR, EVOLUIR E INTEGRAR

## (DO TICKET AO CONHECIMENTO + CHATBOT 2.0)



### OPERAÇÃO CONTÍNUA E GOVERNANÇA

Depois do piloto, o chatbot entra na rotina do atendimento. A qualidade deixa de ser um evento e passa a ser um processo. Isso começa por papéis claros.

O dono da Base de Conhecimento responde pelo conteúdo: decide o que entra, aprova versões e garante que artigos críticos estejam atualizados. **A curadoria de IA cuida das regras do jogo: define o limiar de assertividade, mantém as mensagens de “não sei” e de fallback, revisa prompts e avalia amostras de respostas.**

Segurança/LGPD zela pelos dados: aplica mascaramento de PII, controla retenção, acessos e logs. A operação coordena filas, monitora SLAs e resolve gargalos do dia a dia. **Quando cada um sabe o seu papel, o bot melhora sem improviso.**

Os guardrails técnicos sustentam a segurança. Toda resposta automática precisa estar ancorada em trechos da Base de Conhecimento e registrar quais foram usados.

**O sistema só publica quando a assertividade está acima do limiar combinado; abaixo disso, o bot assume “não sei” e segue o fallback definido (abrir ticket, transferir para humano ou pedir mais dados).**

Permissões determinam quem pode editar conteúdo, alterar políticas e acionar integrações; os logs guardam o histórico de decisões e são auditáveis. **Para temas sensíveis (financeiro, credenciais, dados pessoais), a política pode exigir aprovação humana antes de qualquer ação.**

A manutenção é semanal e leve. Reserve um encontro rápido para revisar uma amostra de conversas reais, confirmar o tom e a aderência à Base e coletar oportunidades de melhoria.

Tudo que o bot não soube responder vira item de backlog: primeiro ajustamos a pergunta (evitando ambiguidades), depois criamos ou melhoramos o artigo correspondente.



# OPERAR, EVOLUIR E INTEGRAR

## (DO TICKET AO CONHECIMENTO + CHATBOT 2.0)



### Esse é o ciclo “ticket → conhecimento”:

cada dúvida recorrente vira conteúdo oficial, aumentando a cobertura e reduzindo a necessidade de escalonamento.

**Defina prazos simples** (por exemplo, publicar artigos prioritários em até cinco dias) e indique responsáveis por tema para que nada se perca.

**As métricas guiam as decisões sem complicar.** Acompanhamos assertividade média, percentual resolvido pelo bot, motivos de fallback, MTTR, FCR, CSAT e tendência de reaberturas. Quando algo sair do normal (ex: queda de assertividade, aumento de fallback ou de reclamações) pausamos a expansão, ajustamos a Base de Conhecimento ou as mensagens do bot e só então retomamos.

**Transparência com o time ajuda:** mostre o painel, explique o porquê das mudanças e reconheça quem contribuiu com artigos e correções que reduziram o retrabalho.

**Por fim, documente acordos de qualidade fáceis de entender:** zero “alucinação crítica”, respostas automáticas sempre rastreáveis, respeito à LGPD e melhora contínua a cada quinzena. Com esse jeito de operar com papéis definidos, guardrails ativos, rituais curtos e métricas à vista, você mantém o serviço seguro, previsível e em evolução.

### DO TICKET AO CONHECIMENTO

Toda dúvida recorrente é um sinal de que falta conteúdo explícito. Por isso, trate tickets complexos ou repetidos como matéria-prima para novos artigos da Base de Conhecimento.

### Priorize pelo impacto e volume (80/20):

comece pelas questões que mais consomem tempo ou geram reabertura. Assim, cada esforço de escrita devolve ganho imediato no atendimento.





# OPERAR, EVOLUIR E INTEGRAR

## (DO TICKET AO CONHECIMENTO + CHATBOT 2.0)



**O fluxo é simples e sempre igual:** identificar → redigir → revisar → publicar. Na identificação, escolha um ticket real e registre o motivo de ter virado conteúdo (volume, risco, custo e etc). Na redação, use o template padrão: quando usar, pré-requisitos, passo a passo, exceções e “definition of done”. Na revisão, faça a conferência de precisão, LGPD/PII e clareza de linguagem. Na publicação, defina dono do artigo, versão, data de revisão e trechos (chunks) com títulos objetivos para facilitar a recuperação pela IA.

**Relacione sempre o artigo ao ticket de origem.** Esse vínculo permite auditar por que o conteúdo existe, medir a queda de reaberturas e comprovar ganho de tempo na próxima ocorrência do mesmo problema. Se o ticket vier de um cliente estratégico, anexe também observações contextuais (ambiente, restrições), mantendo o conteúdo genérico e sem dados pessoais.

**O feedback fecha o ciclo.** Use as avaliações do cliente (“resolveu / não resolveu”) e os comentários dos agentes para refinar passos, exemplos e mensagens. Quando uma resposta automática aciona fallback, verifique quais trechos foram usados e ajuste título, sinônimos e nível de detalhe. Pequenos ajustes de linguagem, como trocar jargão por termos do glossário, costumam aumentar a assertividade sem reescrever tudo.

Com essa rotina, cada novo artigo aumenta a cobertura, melhora a assertividade da resposta e reduz a necessidade de escalonamento.

### **Em poucas semanas, você percebe três efeitos:**

menos tempo gasto explicando o básico, mais consistência entre analistas e um bot cada vez mais autônomo e confiável, porque aprende com o que realmente acontece no seu atendimento.



# OPERAR, EVOLUIR E INTEGRAR

## (DO TICKET AO CONHECIMENTO + CHATBOT 2.0)



Na prática, o **Chatbot 2.0 funciona de forma orquestrada**. Cada evento da conversa ou do ticket (criado, classificado, atualizado) pode acionar um workflow em uma ferramenta low-code (como n8n) ou uma API interna.

**Para tarefas rápidas**, usa-se uma chamada síncrona (o sistema responde na hora). **Para tarefas mais longas**, o fluxo roda assíncrono e o chatbot acompanha por callback ou consulta periódica, registrando tudo na linha do tempo do atendimento. Assim, “intenção → ação” vira um caminho curto e rastreável.

### Exemplos práticos (intenção → ação):

“Segunda via de boleto” → consultar faturas abertas → gerar/linkar PDF → devolver protocolo/URL ao cliente.

“Confirmar pagamento” → checar gateway/ERP → atualizar status do ticket e orientar próximos passos.

“Problema de acesso” → verificar status de serviço/licença → reset padronizado (com aprovação se sensível).

“Queda de rede” → executar checagens técnicas (ping/serviço) → abrir incidente com categoria correta se persistir.

A interoperabilidade precisa ser segura e resiliente. Os fluxos operam com permissões mínimas, logs/auditoria, limites de tempo (timeouts) e retries controlados; ações sensíveis pedem aprovação humana.

**Se algo falhar** (API fora, credencial inválida), o bot assume o limite, explica o que ocorreu e aplica fallback (Plano B): abre um chamado completo para o time, com o contexto do que tentou fazer.

**Para ativar integrações com segurança**, use feature flags (chaves de ativação) para liberar aos poucos, começando fora do horário comercial ou para uma pequena parte do tráfego, e aplique rate limits (limites de chamadas) para não sobrecarregar os sistemas.

O ganho vem de encurtar o caminho entre a intenção e o desfecho. Em vez de repetir instruções, o Chatbot 2.0 faz o trabalho: consulta, atualiza, agenda, valida, confirma — sempre dentro das regras, com transparência e LGPD em dia. Medindo taxa de sucesso das ações, tempo até conclusão e motivos de erro, você expande o escopo com segurança, de processos administrativos simples para casos técnicos mais elaborados.



# OPERAR, EVOLUIR E INTEGRAR

## (DO TICKET AO CONHECIMENTO + CHATBOT 2.0)

06

### Ser IA ou Não.

Dizer (ou não) que a resposta vem de uma IA é uma escolha de experiência e confiança. As duas formas podem funcionar, o ponto é ser claro com a pessoa do outro lado.

Se você assume que é IA, fica mais fácil alinhar expectativas: o usuário entende que o assistente pode não saber tudo e aceita melhor quando for preciso abrir um chamado ou chamar um atendente. Em alguns segmentos, inclusive, informar isso é uma exigência de regras internas ou do próprio mercado.

Se você não destaca que é IA, a conversa pode ficar mais leve em tarefas rápidas (segunda via, status, agendamento). Mas a honestidade continua valendo: se o assistente não tiver certeza, é melhor dizer isso e oferecer o próximo passo (abrir chamado ou passar para uma pessoa), sem rodeios.

### Caminho prático:

escolha a abordagem por canal e por momento (site, WhatsApp, área logada) e teste o texto. O que importa não é “parecer humano”, e sim resolver com segurança: dizer o que sabe, assumir quando não sabe e indicar o caminho para a solução.



# COMO O MILVUS APLICA IA

## NA PRÁTICA PARA TRANSFORMAR O SERVICE DESK



Mesmo que você ainda não utilize o **Milvus**,

**é assim que aplicamos IA do jeito certo:**

organizamos a base de conhecimento, aceleramos a abertura e a classificação de chamados, sugerimos respostas para os analistas e sustentamos um ciclo contínuo de melhoria. O foco é resultado prático com governança, não “IA mágica”.

No **Milvus**, é possível habilitar análise de sentimento, sumarização e rascunhos de abertura/encerramento para apoiar o analista antes de qualquer exposição ao cliente. A operação pode adotar um limiar de assertividade (definido por você) para decidir o que vai a público e, quando a base não cobre, o fluxo recomendado é assumir “não sei” e seguir para fallback. A base de conhecimento pode ser estruturada e tokenizada para recuperar o trecho certo e reduzir custo de processamento.

O **Milvus** suporta orquestração low-code (ex.: n8n) e integrações via API/webhook: a cada movimento do ticket (criação, classificação, atualização), é possível disparar ações em sistemas do negócio (como emitir segunda via, confirmar pagamento ou consultar status) com registro dos eventos no histórico do ticket. Esses fluxos são opcionais e dependem de configuração e das permissões do seu ambiente.

**Contamos também com uma integração direta com o ChatGPT (OpenAI),**  
o que potencializa a gestão de chamados, comentários e finalização.

**Solicite uma demonstração gratuita**



# COMO O MILVUS APLICA IA

## NA PRÁTICA PARA TRANSFORMAR O SERVICE DESK

07

### Aqui, a IA pode ser utilizada para:

#### Gestão de Tickets

- **Sentimento do Ticket:** O Milvus analisa o sentimento do ticket para identificar a percepção do cliente, ajudando a priorizar e personalizar o atendimento.
- **Sugerir texto para Início do Atendimento:** A IA sugere textos prontos para iniciar o atendimento, agilizando a resposta inicial e padronizando a comunicação.
- **Questões para ajudar a entender melhor a solicitação:** O ChatGPT pode gerar perguntas relevantes para que o atendente obtenha mais detalhes do cliente, otimizando o entendimento da solicitação.
- **Texto para Finalizar o Ticket:** A IA auxilia na criação de textos claros e completos para finalizar o ticket, garantindo que todas as informações importantes sejam registradas.
- **Criar novos comentários:** Permite gerar comentários internos ou externos diretamente nos tickets, seja com base na inteligência do ChatGPT ou utilizando informações da base de conhecimento.
- **Criar novas descrições de E-mails:** A IA pode criar descrições de e-mails para notificações e comunicações relacionadas aos tickets, utilizando o ChatGPT ou a base de conhecimento.





# COMO O MILVUS APLICA IA

## NA PRÁTICA PARA TRANSFORMAR O SERVICE DESK

07

### Automação e Otimização

- ➔ **Powershell e Scripts:** A inteligência artificial pode auxiliar na geração de comandos Powershell e scripts, automatizando tarefas e solucionando problemas com mais agilidade.
- ➔ **Transcrever áudio:** O Milvus pode transcrever áudios (como chamadas ou mensagens de voz) para texto, facilitando o registro e a consulta de informações nos tickets.
- ➔ **Melhorar e corrigir a ortografia e gramática do texto em comentário:** Garante que os textos dos comentários estejam corretos e padronizados, melhorando a comunicação interna e externa.

### Base de Conhecimento

- ➔ **Tokenizar a base de conhecimento:** A IA ajuda a estruturar e categorizar a base de conhecimento, tornando-a mais eficiente para buscas e sugestões de soluções específicas.
- ➔ **Sugerir uma Solução para o Ticket:** A inteligência artificial pode sugerir soluções para os tickets, buscando informações tanto no ChatGPT quanto na base de conhecimento já existente.

### Respostas e Interações

- ➔ **Criar resposta para chat:** A IA pode gerar respostas automáticas para o chat, usando tanto o ChatGPT para respostas mais dinâmicas quanto a base de conhecimento para informações padronizadas.



# **A.CHECKLIST DE SUCESSO**

(1 PÁGINA)

## **1) Base de Conhecimento**

- ❑ Dono por tema definido (quem cria/revisa/arquiva).
- ❑ Artigos homologados, com versão e data de revisão/expiração.
- ❑ Conteúdo padronizado (passo a passo, exceções, definição de “pronto”).
- ❑ Tokenização em trechos curtos (chunks) e títulos claros.
- ❑ Rastreabilidade: resposta do bot cita os trechos usados.

## **2) Segurança & LGPD**

- ❑ PII (dados pessoais) mascarada; prints limpos.
- ❑ Permissões por perfil; logs/auditoria ativos.
- ❑ Política de retenção de dados definida.
- ❑ Nenhum incidente de LGPD pendente.

## **3) Política de Resposta Segura**

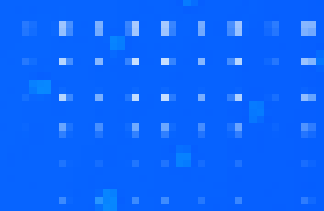
- ❑ Regra de ouro: bot só responde com conteúdo homologado. “Não sei” configurado (mensagem curta e clara).
- ❑ Fallback pronto: abrir ticket, transferir para humano ou pedir dados.
- ❑ Desambiguação (sugerir reformulações) ativa.
- ❑ Limiar de assertividade definido (ex.: 70%) e aplicado.

## **4) Piloto & Operação**

- ❑ Piloto off-hours rodando há  $\geq 2$  semanas sem “alucinação crítica”.
- ❑ Playbooks de mão dupla: quando responder / quando transferir.
- ❑ Plano de contingência documentado (pausar N1, manter abertura/classificação).
- ❑ Revisão semanal de 15 min (QA por amostra) agendada.

## **5) Métricas (linha de base vs. piloto)**

- ❑ % resolvidos pelo bot subindo.
- ❑ MTTR (tempo médio para resolver) caindo.
- ❑ FCR (resolução no primeiro contato) estável/subindo.
- ❑ CSAT (satisfação)  $\geq$  linha de base.
- ❑ Taxa de fallback controlada; alucinação crítica = 0.
- ❑ Cobertura da base nas intenções do piloto  $\geq 80\%$ .



**milvus**.com.br

