

IT a anatomie firmy

(Big Data ve strojírenské firmě)

(pracovní dokument)



Markéta Mánková

VŠE Praha, 2026



Mapa dokumentu podle kapitol textu (s odkazy)

[1] Big Data v kontextu výroby	
[2] Možnosti ukládání Big Data	
[3] Optimalizace řešení pro skalární data	[4] Optimalizace řešení pro křivková data
[5] Aplikace pro skalární data	[6] Aplikace pro křivková data
[7] Vyhodnocení přínosů	



Účelem dokumentu je prezentovat metody **optimalizace zpracování velkých výrobních dat** s důrazem na dva aspekty – snížení objemu přenášených a ukládaných dat a omezení nákladů spojených s reklamami ve výrobním procesu.

Obsah

1.	<i>Big Data v kontextu výroby</i>	5
1.1	Definice a charakteristika Big Data	5
1.2	Oblasti generování dat ve výrobním prostředí	5
1.2.1	Klasifikace dle oblasti vzniku a aplikace dat	5
1.2.2	Klasifikace dle typu zpracování	6
1.3	Výzvy spojené s Big Data	6
1.4	Výhody implementace Big Data	7
1.5	Závěry	7
2.	<i>Možnosti ukládání Big data</i>	8
2.1	Technologie pro ukládání	8
2.1.1	Distribuované souborové systémy	8
2.1.2	NoSQL databáze	8
2.1.3	NewSQL databáze	9
2.1.4	Dotazovací platformy pro Big Data	9
2.2	Přístupy k ukládání	9
2.2.1	On-Premise infrastruktura	9
2.2.2	Cloud computing	9
2.3	Závěry	10
3.	<i>Optimalizační řešení pro skalární data</i>	11
3.1	Předzpracování dat	11
3.2	Informační hodnota parametrů	11
3.3	Detekce anomálií	12
3.4	Závěry	14
4.	<i>Optimalizační řešení pro křivková data</i>	15
4.1	Předzpracování dat	15
4.2	Komprese	15
4.3	Analýza významnosti	16
4.4	Detekce anomálií	16
4.5	Závěry	18
5.	<i>Aplikace navržených metod pro skalární data</i>	19
5.1	Předzpracování dat	19
5.2	Selekce parametrů a jejich popis	19
5.3	Detekce anomálií na skalárech	23
5.4	Závěry	26
6.	<i>Aplikace navržených metod pro křivková data</i>	27
6.1	Předzpracování a čištění dat	27
6.2	Komprese procesních křivek	29

6.3	Významnost jednotlivých úseků křivky.....	35
6.4	Detekce anomálií na křivkách.....	37
6.5	Závěry	50
7.	<i>Vyhodnocení přínosů</i>	51
7.1	Detekce anomálií na skalárních datech.....	51
7.2	Komprese křivkových dat	51
7.3	Detekce anomálií na křivkových datech	52
7.4	Závěry	53
8.	<i>Závěr</i>	54
	<i>Zdroje</i>	55

1. Big Data v kontextu výroby



Účelem je vymezení velkých objemů dat, resp. Big Data s orientací na jejich řešení a užití ve strojírenských firmách.

1.1 Definice a charakteristika Big Data

Pojem Big Data označuje masivní a komplexní datové sady, které **svými nároky značně přesahují možnosti tradičních systémů** pro správu dat. Zatímco tradiční data jsou primárně strukturovaná a organizovaná v relačních databázích, Big Data **zahrnují data v různých formátech**, jejichž zpracování vyžaduje pokročilé analytické přístupy a distribuované systémy schopné efektivně pracovat v obrovském měřítku (Badman & Kosinski, 2024).

Specifika, která činí velká data unikátními oproti tradičním datovým sadám, jsou **definována pěti charakteristikami, často označovanými jako „V“** Big Data: Objem (Volume), Rychlost (Velocity), Rozmanitost (Variety), Věrohodnost (Veracity) a Hodnota (Value) (Badman & Kosinski, 2024).

Prvním aspektem Big Data je objem, který se vztahuje k **množství dat**, jež jsou dnes generována a ukládána. U rozsáhlých datových sad se běžně pracuje s velmi vysokými počty záznamů, na které tradiční úložné systémy často nestačí. Právě omezená škálovatelnost klasických řešení je jedním z důvodů, proč se objevuje potřeba nových přístupů k ukládání a správě dat (Badman & Kosinski, 2024).

Druhou charakteristikou je **rychlost, tedy tempo, jakým data do systému přicházejí a jak rychle je nutné je zpracovávat**. V současnosti vznikají data téměř nepřetržitě a často v reálném čase – schopnost reagovat na data okamžitě se tak často stává primárním požadavkem moderních informačních systémů (Badman & Kosinski, 2024).

Třetí dimenzí je **rozmanitost, popisující šíři datových formátů**, se kterými je potřeba pracovat. Vedle klasických strukturovaných dat se stále častěji objevují také data polostrukturovaná a nestrukturovaná, která nemají pevně dané schéma, a to pochopitelně zvyšuje nároky na jejich ukládání i integraci (Badman & Kosinski, 2024).

Čtvrtým pilířem je **věrohodnost, tedy kvalita a spolehlivost dat** – data pocházejí z různých zdrojů a mohou obsahovat chyby. Pokud nejsou tyto problémy řešeny, může dojít ke zkresleným výstupům a chybným rozhodnutím, a je tedy nutné zavádět procesy čištění a validace dat, jež pomáhají zajišťovat jejich správnost (Badman & Kosinski, 2024).

Posledním aspektem modelu 5V je **hodnota, která vyjadřuje skutečný přínos dat pro organizaci**; samotné shromažďování dat nemá smysl, pokud data nejsou dále analyzována a využita (Badman & Kosinski, 2024).

1.2 Oblasti generování dat ve výrobním prostředí

Výrobní průmysl neustále prochází revolucí řízenou daty, která transformuje tradiční provozy na vysoce optimalizované prostředí tzv. chytré výroby. V těchto prostředích **dochází k bezprecedentní produkci dat, která pocházejí ze starších automatizačních systémů a senzorových sítí**, stejně jako z nově nastupujících technologií. Nízkoúrovňová granulózní data jsou zaznamenávána v reálném čase a slouží k vytváření „výrobní inteligence“, která podporuje přesné a včasné rozhodování (O'Donovan et al., 2015).

Dle systematické studie O'Donovana et al. (2015) lze data a výzkum s nimi spojený **rozdělit primárně podle oblastí výroby, kde data vznikají, a podle typu analytiky**, která je k jejich zpracování využita.

1.2.1 Klasifikace dle oblastí vzniku a aplikace dat

Autoři zmiňují **několik konkrétních oblastí výrobního prostředí**, které generují **specifická data** a vyžadují tak **rozdílné přístupy zpracování**:

- Data v oblasti procesů a plánování, která slouží primárně k optimalizaci výrobních postupů a souvisí s tokem výroby a samotnou efektivitou operací,
- data v oblasti údržby a diagnostiky sbíraná za účelem prediktivní údržby a diagnostiky v reálném čase,

- data z kategorie řízení kvality, díky kterým lze zajistit sledování a následně dostatečnou výrobní kvalitu,
- podniková data, která procházejí napříč celou organizací a slouží k podpoře rozhodování na vyšší úrovni řízení podniku,
- data využívaná pro design a virtuální výrobu,
- specifická data zaměřená na plánování aktivit a přímé řízení procesů ve výrobním prostředí,
- data monitorující faktory související s environmentálními dopady, bezpečností práce, a tedy i ochranou zdraví (O'Donovan et al., 2015).

1.2.2 Klasifikace dle typu zpracování

Pro efektivní využití velkých objemů dat lze **klasifikovat způsob analýzy**, kterým jsou data zpracovávána, a O'Donovan et al. (2015) ji klasifikují **do tří úrovní pokročilosti**:

- Deskriptivní analytika, která se zaměřuje na popis struktury, vztahů a významu dat,
- prediktivní analytika, jenž využívá dostupná data k předpovídání budoucích výsledků
- preskriptivní analytika – nejpokročilejší typ analytiky, který na základě dostupných dat přímo předepisuje konkrétní akce, které by měly být provedeny v případě různých událostí.

Data ve výrobním prostředí nevznikají jednotně, ale v různých částech procesu a s různým účelem; pro další postup jsou **podstatná především data přímo spojená s průběhem výroby a kontrolou kvality**, protože právě ta nesou informaci o tom, zda se výrobní kus chová standardně, nebo dochází k odchylkám.

1.3 Výzvy spojené s Big Data

Správa a potenciálně i analýza velkých datových sad představují řadu výzev různých druhů, jenž dokážou hybat s efektivitou a úspěšností implementace těchto technologií v organizaci.

Obecně v organizacích je i v dnešní digitální době nejproblematictější **nedostatečné porozumění problematice a přijetí konceptu velkých dat** v organizacích; spousta podniků si stále neuvědomuje potenciál využívání velkých dat a zároveň podceňuje náročnost jejich správy, čímž eventuálně může docházet k neefektivnímu plánování IT infrastruktury v samém počátku (Naeem et al., 2022).

Další výzvou je dle Naeema et al. (2022) **změť dostupných technologií a nástrojů** pro práci s velkými daty – vzhledem k rozmanitému ekosystému Big Data platform, různých databází, analytických nástrojů a nepřehledného množství cloudových služeb je pro organizace, které o implementaci Big Data uvažují a nemají dostatek interních expertů, velmi obtížné zvolit správné technologické řešení.

Významnou překážkou je také **samotná analytická náročnost**, která klade vysoké nároky na lidské zdroje, protože proces transformace surových dat na užitečné znalosti pro strategické rozhodování vyžaduje multidisciplinární dovednosti. Dle Naeema et al. (2022) nestačí pouze data shromáždit; organizace potřebují vysoce kvalifikované týmy datových inženýrů, kteří připraví data tak, aby je následně mohli zpracovat datoví vědci. Nedostatek takto specifických expertů na trhu, kteří by zvládli práci s nestrukturovanou analytikou, představuje pro mnoho firem kritické úzké hrdlo.

S rostoucím objemem dat roste přímo úměrně i **riziko spjaté s bezpečností a ochranou soukromí** – správa dat v takovém prostředí vyžaduje sofistikované mechanismy ochrany, neboť zpravidla obsahuje přinejmenším data podléhající ochraně osobních údajů. Nedostatečné zabezpečení může nejen v případě útoku vést k bezpečnostním incidentům v různých mírách; situaci dále komplikuje skutečnost, že mnoho organizací využívá cloudové služby třetích stran kvůli nedostatečným zdrojům pro vlastní infrastrukturu a tím je obecně vzato zvýšeno riziko neautorizovaného přístupu k datům (Naeem et al., 2022).

Faktor, který je nutné zvážit při úvaze o zpracovávání Big Data, je **proces ETL (Extract, Transform, Load)**. Takové organizace, které disponují vlastním systémem, čelí dle definice procesu transformace, čištění od redundancí a chyb a následnému uložení dat do datového skladu v optimalizovaném formátu – takový proces je každopádně časově i výpočetně náročný a ve společnosti může zabírat až 50 % celkového času potřebného pro datovou analytiku (Naeem et al., 2022).

V neposlední řadě je třeba zmínit **potřebu škálovatelnosti fyzické infrastruktury** díky skutečnosti, že tradiční IT infrastruktura není navržena pro práci s nepředvídatelným a extrémně velkým objemem

dat. Podniky tedy musí implementovat distribuované systémy, které jsou schopny dynamické škálovatelnosti aktuálních potřeb zpracování (Naeem et al., 2022).

Z uvedených výzev vyplývá potřeba **navrhovat řešení, která reflektují jak rostoucí objem dat a s tím spojené nároky na infrastrukturu, tak i jejich analytickou náročnost**. V daném kontextu se proto jedná především o snahu efektivně redukovat objem ukládaných dat s cílem omezení nároků na datovou infrastrukturu a zároveň automatizovat odhalování nestandardních průběhů bez nutnosti manuálního vyhodnocování – nebo pouze s minimálním zásahem.

1.4 Výhody implementace Big Data

Jedním z hlavních **přínosů Big Data** je schopnost **detailně sledovat a analyzovat průběh výrobních procesů v reálném čase, díky které lze přejít z reaktivního řízení jakosti k proaktivnímu**. Sběrem kombinace dat ze senzorů, řídicích systémů a z dříve uložených provozních záznamů lze odhalit poměrně přesný obraz toho, jakým způsobem výroba funguje a jak se chová. Za předpokladu, že jsou tato data průběžně vyhodnocována, si lze poměrně rychle všimnout situací, kdy se začne proces odchýlovat od obvyklého nastavení. Díky tomu je **možné zasáhnout ještě před tím, než se problém naplno objeví** – v praxi se to projevuje menším počtem vadných kusů a obecně vyrovnanějším průběhem výroby (O'Donovan et al., 2015).

Big Data rovněž **podporují optimalizaci výrobních operací z hlediska výkonnosti a dostupnosti** zařízení; analýza reálných provozních dat totiž umožňuje výrazně lepší odhad pro plánování kapacit, minimalizaci prostojových časů, a tedy i optimalizaci výrobní sekvence. O'Donovan et al. (2015) poukazují na skutečnost, že propojení dat z různých úrovní řízení vytváří celistvý obraz o průběhu výrobních toků, díky němuž spíše lze kvalifikovaně rozhodovat v rámci celého podniku.

Implementace Big Data technologií **se promítá také do finančních ukazatelů výroby**, a to skrze snížení míry zmetkovitosti a nižší počet reklamací či například efektivnějším plánováním údržby. Jak uvádějí O'Donovan et al. (2015), data-driven přístup napomáhá přesnější alokaci zdrojů a eliminuje rozhodování, jež by jinak vycházelo pouze z odhadů či zkušenosti jednotlivců.

Přínosy implementace Big Data technologií jsou tedy **v průmyslovém prostředí konkrétní a měřitelné** – od snížení počtu vadných kusů a reklamací až po efektivnější využití výrobních kapacit a podpůrných zdrojů. V kontextu analyzovaného výrobního závodu představuje právě schopnost včasné detekce odchylek a proaktivního řízení kvality důležitý argument pro rozvoj analytických přístupů nad dostupnými výrobními daty.

1.5 Závěry



- Problematika Big Data ve výrobním prostředí není omezena pouze na schopnost data ukládat, ale **zahrnuje celý řetězec rozhodnutí od jejich sběru přes správu, předzpracování a analytické využití až po vyhodnocení** jejich skutečné hodnoty pro podnik.
- V kontextu výroby jsou podstatná zejména **data vznikající přímo v průběhu výrobního procesu a při kontrole kvality**, protože právě tato data umožňují sledovat stabilitu procesu, identifikovat odchylky a vytvářet podklady pro proaktivní řízení jakosti.

2. Možnosti ukládání Big data



Účelem je charakterizovat různé technologické možnosti, jejich efekty a případná omezení spojené s ukládáním velkých objemů dat.

Ukládání velkých dat představuje **fundamentální vrstvu datově orientovaných architektur** a tvoří předpoklad pro efektivní zpracování a následnou monetizaci dat. Technologie na ukládání, a tedy v určitém smyslu i uchovávání Big Data jsou navrhovány s cílem **zajistit dlouhodobě udržitelnou správu dat v podmínkách exponenciálního růstu jejich objemu** a dalších již zmíněných parametrů. Ideální systém pro ukládání velkých dat by podle Strohbacha et al. (2016) měl splňovat několik často protichůdných požadavků; kromě prakticky neomezené škálovatelnosti objemu dat a schopnosti **zvládat vysoké rychlosti** zápisu i čtení by měl rovněž **flexibilně podporovat různé datové modely** a současně umožňovat práci se strukturovanými, nestrukturovanými nebo polostrukturovanými daty. Z hlediska ochrany soukromí by pak měl ideálně pracovat **výhradně se šifrovanými daty**. Autoři nicméně upozorňují, že všechny tyto požadavky nelze v plném rozsahu současně naplnit, a současné systémy proto představují různé kompromisy mezi výkonem, škálovatelností, konzistencí a bezpečností.

Takové požadavky zásadním způsobem překračují možnosti tradičních relačních databázových systémů, jejichž architektura je historicky založena především na vertikálním škálování a silných transakčních garancích. V reakci na tato omezení dochází k posunu směrem **k distribuovaným architekturám typu shared-nothing**, které preferují horizontální škálování prostřednictvím většího počtu relativně nezávislých výpočetních uzlů (Strohbach et al., 2016).

2.1 Technologie pro ukládání

Dynamický rozvoj datové intenzivních aplikací a současné změny v hardwarových paradigmatech vedly ke vzniku celé řady nových úložných technologií, které se vědomě odklánějí od klasického relačního modelu – tyto systémy často upřednostňují škálovatelnost, dostupnost a výkon před striktním dodržováním transakčních vlastností, zejména okamžité konzistence dat. Na základě architektonických a funkčních charakteristik lze **technologie pro ukládání Big Data rozdělit do několika základních kategorií** (Strohbach et al., 2016).

2.1.1 Distribuované souborové systémy

Základní infrastrukturní vrstvu pro ukládání rozsáhlých objemů **nestrukturovaných dat tvoří distribuované souborové systémy**. V této oblasti se de facto jako standard prosadil **Hadoop Distributed File System**, jehož návrh je optimalizován pro spolehlivé ukládání velmi velkých souborů na komoditním hardwaru. Architektura HDFS je uzpůsobena zejména pro scénáře hromadného nahrávání dat a následné dávkové zpracování, přičemž se soustředí na propustnost spíše než na nízkou latenci přístupu k jednotlivým záznamům (Strohbach et al., 2016).

2.1.2 NoSQL databáze

Významnou skupinu představují databázové systémy souhrnně označované jako NoSQL. Takové technologie využívají alternativní datové modely a zpravidla se vzdávají plné podpory transakčních vlastností ACID ve prospěch vyšší škálovatelnosti a dostupnosti. **V závislosti na použitém datovém modelu lze NoSQL databáze dále členit na několik podkategorií:**

- **Úložiště klíč–hodnota**, která umožňují ukládání dat bez předem definovaného schématu, přičemž přístup k objektům je realizován výhradně prostřednictvím unikátního klíče,
- **sloupcově orientovaná úložiště**, jež organizují data po sloupcích namísto řádků, čímž dosahují vysoké efektivity při kompresi a analytickém zpracování řídkých dat. Data jsou typicky identifikována kombinací klíče řádku, klíče sloupce a časového razítka,
- **dokumentové databáze**, určené pro ukládání polostrukturovaných dat ve formátech typu XML nebo JSON, které umožňují dotazování nad vnitřní strukturou dokumentů bez nutnosti jednotného globálního schématu,

- **grafové databáze**, optimalizované pro práci s vysoce propojenými daty, kde jsou klíčovým prvkem vztahy mezi entitami. Distribuované škálování těchto systémů je však technologicky náročné vzhledem ke komplexitě grafových struktur (Strohbach et al., 2016).

2.1.3 NewSQL databáze

Jako reakce na kompromisy typické pro NoSQL systémy se objevují tzv. NewSQL databáze, jejichž cílem je **propojit výhody horizontální škálovatelnosti s plnou podporou transakčních vlastností** a dotazovacího jazyka SQL. Systémy jsou navrhovány pro provoz v architekturách typu shared-nothing a usilují o dosažení vysokého výkonu na úrovni jednotlivých uzlů při zachování konzistence dat (Strohbach et al., 2016).

2.1.4 Dotazovací platformy pro Big Data

Samostatnou kategorií tvoří dotazovací platformy, které fungují jako logické zprostředkující vrstvy mezi uživatelem a distribuovanými úložišti a které dávají uživatelům možnost analyzovat data uložená v distribuovaných souborových systémech či NoSQL databázích prostřednictvím dotazovacích jazyků inspirovaných SQL. **Základním principem je zde přístup schema-on-read**, kdy je struktura dat definována až v okamžiku dotazu, nikoliv při jejich ukládání. Zmíněný koncept zvyšuje flexibilitu práce s daty v raných fázích jejich životního cyklu (Strohbach et al., 2016).

2.2 Přístupy k ukládání

Rozhodnutí o způsobu organizace ukládání a zpracovávání velkých dat má dlouhodobé **dopady na nákladovou strukturu a řízení rizik** firmy. Ali et al. (2024) poukazují na to, že se podniky dnes typicky pohybují mezi dvěma základními architektonickými přístupy: tradiční lokální infrastrukturou a cloudovým prostředím.

2.2.1 On-Premise infrastruktura

Model on-premise je založen na vlastnictví a správě veškeré infrastruktury přímo organizací – ta zajišťuje provoz datových center, stejně jako serverů, úložišť i podpůrných systémů samostatně. Hlavním argumentem **ve prospěch tohoto řešení je plná kontrola nad daty a jejich fyzickým umístěním**, což lze považovat za relevantní třeba pro subjekty působící v regulovaných odvětvích, zejména ve finančním sektoru nebo ve veřejné správě (Ali et al., 2024).

Lokální přístup je na druhou stranu **spojen s vysokými kapitálovými výdaji a značnými nároky na provozní správu**; organizace musí nést odpovědnost za rozsáhlé množství proměnných – údržba hardwaru, energetické zabezpečení, chlazení, zálohování i obnova po havárii leží v takové situaci na bedrech společnosti. Z pohledu Big Data projektů je problematická především omezená škálovatelnost, neboť rozšíření kapacity vyžaduje časově i finančně náročné investice do infrastruktury nové, a to zase může být v rozporu s dynamikou datově orientovaných aplikací (Ali et al., 2024).

2.2.2 Cloud computing

Cloudové prostředí představuje alternativu založenou na poskytování výpočetních a úložných zdrojů formou služby – organizace tak získávají **přístup k infrastruktuře prostřednictvím internetu a hradí pouze skutečně využitou kapacitu**, čímž dochází k transformaci kapitálových výdajů na provozní (Ali et al., 2024).

Z hlediska implementace lze rozlišit **veřejný, privátní a hybridní model**. Veřejný cloud je charakteristický sdílenou infrastrukturou spravovanou poskytovatelem a jeho výhodou je bezesporně vysoká míra škálovatelnosti a ekonomická efektivita. Privátní cloud je vyhrazen jedné organizaci a nabízí vyšší míru kontroly, avšak za cenu vyšších nákladů. Hybridní model kombinuje oba přístupy a umožňuje diferencované nakládání s daty podle jejich citlivosti a provozních požadavků (Ali et al., 2024).

Pro oblast **Big Data je typický zejména model Infrastructure-as-a-Service**, který poskytuje virtualizované výpočetní zdroje a úložiště. Ali et al. (2024) uvádějí, že většina analyzovaných organizací preferuje cloudové řešení před lokální infrastrukturou, a to zejména z důvodu škálovatelnosti a nižších vstupních nákladů. Řešení jsou navržena tak, aby umožňovala efektivní ukládání a zpracování

rozsáhlých objemů nestrukturovaných dat; tato vlastnost patří mezi základní předpoklady fungování aplikací pracujících s Big Data.

Z uvedeného přehledu vyplývá, že existuje široké spektrum technologických a architektonických přístupů k ukládání velkých dat, které se liší zejména způsobem škálování a mírou konzistence či výkonu.

2.3 Závěry



- **Rostoucí objem výrobních dat** přináší nejen analytické příležitosti, ale také značné nároky na datovou infrastrukturu, výpočetní kapacity, správu dat a jejich dlouhodobé uchovávání.
- Přehled technologií a přístupů k ukládání ukazuje, že moderní Big Data architektury **dokážou škálovat a pracovat s různorodými formáty dat**, avšak každé řešení je spojeno s určitými kompromisy.
- Z toho vyplývá, že samotné rozšiřování úložišť není dlouhodobě dostatečnou odpovědí na růst datových objemů; vhodnější je **hledat také způsoby, jak redukovat objem ukládaných dat** bez ztráty jejich diagnostické hodnoty.
- Analyzované příklady z praxe velkých technologických a průmyslových společností zároveň potvrzují, že **optimalizace práce s daty může mít různé podoby** – od sjednocení datových platforem přes cloudovou škálovatelnost až po snižování energetické a materiálové náročnosti datové infrastruktury.

3. Optimalizační řešení pro skalární data



Účelem je:

- ukázat **možnost zmenšení objemu přenášených a ukládaných dat prostřednictvím vhodné metody**, která umožní efektivnější práci se záznamy bez ztráty jejich diagnostické hodnoty
- **snížení nákladů spojených s výskytem anomálií** – tedy odchylek v průběhu výrobního procesu, které nejsou zachytitelné běžnými monitorovacími technikami.

Optimalizační přístup bude usilovat o **vytvoření prediktivního rámce**, jenž umožní takové anomálie včas rozpoznat a zabránit jejich dopadům, zejména ve formě reklamací nebo zvýšených nákladů na kontrolu kvality.

Pro účely návrhu a ověření řešení byly zvoleny **dvě pilotní výrobní linky** z důvodu dostupnosti kvalitních historických dat – jak křivkových, tak skalárních – a zároveň reprezentativnosti z hlediska charakteru výrobních operací. Přístup je koncipován tak, aby jej bylo možné následně pouze s minimálními změnami rozšířit i na další linky obdobného typu. Celkem bylo získáno **562 souborů s křivkovými daty** (každý reprezentuje jednu křivku) a **21 tabulek se skalárními záznamy** pro jednotlivé stanice.

Skalární data generovaná v rámci pilotní linky č. 1 zastupují část informací generovaných výrobními linkami a obvykle nesou informaci o parametrech, které popisují průběh různých fyzikálních veličin.

3.1 Předzpracování dat

Plán předzpracování skalárních dat se týká především **zajištění konzistence vstupního datasetu** a eliminace zkreslení způsobené nekvalitními záznamy. V případě asymetrie rozdělení parametrů, která se v nějakém rozsahu očekává, bude aplikována transformační metoda navržená Yeo a Johnsonem (2000). Metoda, která je ve firmě standardem, patří do rodiny mocninných transformací a jejím cílem je zlepšit rozdělení dat nebo jejich přiblížení k normálnímu rozdělení (Yeo & Johnson, 2000).

Data budou **načítána z externího datového souboru ve formátu CSV** a součástí předzpracování bude odstranění duplicitních záznamů, neboť ve výrobní praxi se takové záznamy standardně vyskytují. Pokud se v datech budou vyskytovat vícečetná měření se stejným identifikátorem výrobku, bude zachován pouze nejnovější ze záznamů určený na základě časové značky výsledku měření. Vzhledem k časovému ohraničení dat **budou vyjmuty záznamy, které svoji cestu výrobní linkou začaly před daným časovým intervalem a byly dokončeny v něm** – tedy takové, jejichž identifikátor lze nalézt pouze na navazujících stanicích, nikoliv na stanici č. 1.

Datová sada bude následně **zpracovávána ve dvou navazujících rovinách** – první rovina bude zaměřena na posouzení informační hodnoty jednotlivých parametrů pro rozhodnutí o návrhu zachování či vyloučení v optimalizované datové reprezentaci a druhá rovina bude vybrané informačně hodnotné parametry zpracovávat ve snaze detekovat anomálie.

3.2 Informační hodnota parametrů

Hodnocení informační hodnoty skalárních parametrů má za cíl identifikovat ty veličiny, které přinášejí relevantní informaci o chování výrobního procesu, a naopak **omezit vliv parametrů redundantních nebo analyticky nevýznamných**. Vzhledem k vysokému počtu dostupných skalárních měření a jejich rozdílné kvalitě nebude účelné pracovat se všemi parametry v plném rozsahu, neboť by to mohlo vést ke zvýšené výpočetní náročnosti, a především ke ztížení analýzy.

Při návrhu analytické metodiky byly primárně zvažovány moderní analytické přístupy, které nabízí strukturovaný a kvantifikovatelný základ pro automatizovaný výběr parametrů. Primárně byla analyzována **metoda prahování rozptylu**, jež k hodnocení užitečnosti parametrů využívá hodnotu jejich disperze a vychází z předpokladu, že parametr s vyšším rozptylem je pravděpodobně informativnější, přičemž parametry s rozptylem blížícím se nule lze odfiltrovat (Kamalov et al., 2025).

Aplikace metody na takto specifická výrobní data by **neposkytla dostatečně spolehlivé výsledky** kvůli povaze samotného procesu a dlouhodobé snaze o co nejmenší rozptyl jednotlivých parametrů v rámci optimalizace výroby – i ty nejdůležitější parametry se pohybují v rámci velmi úzkých tolerančních hranic. Pokud by se výběr řídil výpočty dle rozsahu variance, matematický filtr by mylně označil

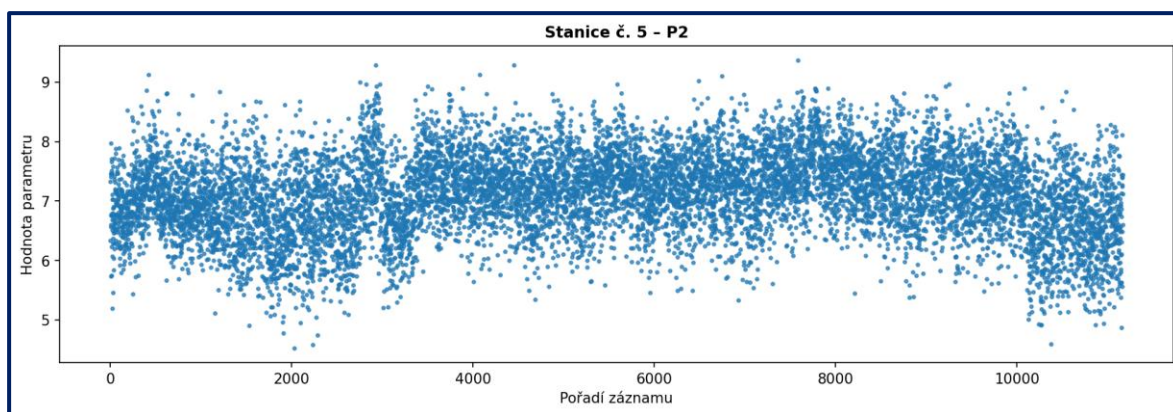
stěžejní výrobní parametry za degenerované či téměř konstantní a z analýzy by je nesprávně vyloučil (Kamalov et al., 2025).

Z tohoto důvodu nebylo možné moderní výpočetní postupy použít. **Selekce parametrů se bude namísto toho opírat výhradně o expertní výběr**, což plně odpovídá tradičním postupům v oblastech se složitými závislostmi. Celý proces bude probíhat skrze expertní konzultace s procesními specialisty, kteří disponují detailní znalostí fyzikální podstaty měřených veličin i technologických omezení jednotlivých výrobních stanic, a mohou tak přesně určit informační hodnotu parametru bez ohledu na jeho aktuální statistickou variabilitu.

Výsledkem této fáze bude **seznam parametrů**, které jsou z hlediska další analýzy relevantní, přičemž surová data zůstanou zachována pro případnou zpětnou dohledatelnost.

3.3 Detekce anomálií

Cílem fáze detekce anomálií je **identifikovat výrobní kusy, které sice formálně splnily toleranční limity** na všech stanicích a byly vyhodnoceny jako dobré, avšak jejich procesní parametry se statisticky vymykají chování běžné produkce. Takové kusy nelze zachytit standardními kontrolními mechanismy založenými na pevně daných mezích, přesto mohou indikovat nestandardní průběh výrobního procesu.



Obrázek 3.1: Ukázka časové řady parametru z pilotní linky č.1 (zdroj: autorka, 2026)

Na obrázku je znázorněn bodový **graf jednoho z parametrů výrobní linky**, který slouží k ilustraci důvodů pro zavedení metody detekce anomálií. Zobrazeny jsou **pouze hodnoty odpovídající kusům označeným jako OK**, přesto je z průběhu patrné, že některé z nich se svým chováním od ostatních odlišují natolik, že by mohly být při bližším posouzení považovány za nestandardní. Na tomto grafu ani na ostatních grafech pracujících se skalárními daty nejsou z důvodu anonymizace uváděny názvy parametrů ani jejich jednotky.

I v rámci výrobků **splňujících stanovená kritéria kvality** je viditelný rozptyl hodnot sledovaného parametru. V určitých úsecích se hodnoty koncentrují v relativně úzkém pásmu, jinde je naopak patrná větší variabilita. Zároveň se vyskytují i jednotlivé body, které z tohoto „typického“ chování vybočují – takové případy však běžné kontrolní mechanismy nezachytí, protože nedochází k překročení předepsaných tolerančních mezí.

Vstupem do této fáze je **redukovaný dataset** obsahující pouze parametry vybrané v předchozí fázi expertní selekcí. Výpočet je prováděn výhradně nad záznamy korektně dokončených výrobků, čímž je zajištěno, že referenční chování odpovídá skutečnému provoznímu standardu výrobní linky.

Původně byla pro kvantifikaci odchylky každého kusu **zamýšlena metoda klasického z-skóre**, jež vyjadřuje odchylku od průměru v jednotkách směrodatné odchylky. Hodnoty, které následně překračují určitou hranici, jsou považovány za anomální. Je nutné zmínit, že takový přístup je poměrně citlivý na přítomnost extrémních hodnot, které se ve výrobním procesu každopádně objevují a potenciálně by tak mohly nepříznivě ovlivnit jak průměr, tak směrodatnou odchylku (Rousseeuw & Hubert, 2018).

Jako **vhodnější varianta se tedy nabízí využití robustní varianty z-skóre**, která nahrazuje klasické statistické charakteristiky jejich robustními protějšky. Konkrétně je průměr nahrazen mediánem a směrodatná odchylka mediánovou absolutní odchylkou (dále jen MAD), která je vůči extrémním hodnotám výrazně méně citlivá. Výpočet probíhá dle vztahu

$$z = \frac{(x - \text{medián}(x))}{MAD(x)}, \quad (3.1)$$

kde z je hodnota robustního z-skóre, x je hodnota měřeného parametru, **medián je středová hodnota distribuce daného parametru napříč všemi kusy a MAD je medián absolutních odchylek od mediánu**. Autoři uvádějí, že pro zajištění konzistence MAD s klasickou směrodatnou odchylkou při normálním rozdělení je hodnota MAD škálována konstantou 1,4826; jelikož je robustní z-skóre v rámci následně normalizováno dělením svou maximální hodnotou, konstanta se ve výpočtu vyruší a její aplikace nemá na výsledné anomaly score žádný vliv, a tak není konstanta ve výpočtu dále uvažována (Rousseeuw & Hubert, 2018).

Pro každý parametr je následně vypočtena absolutní hodnota robustního z-skóre, která je normalizována na interval $(0, 1)$ podělením maximální hodnotou z-skóre v daném parametru dle následující rovnice:

$$z^{norm} = \frac{|z|}{\max(|z|)}, \quad (3.2)$$

kde z^{norm} je normalizované absolutní robustní z-skóre, $|z|$ je absolutní hodnota robustního z-skóre (velikost odchylky bez ohledu na směr) a $\max(|z|)$ je maximální absolutní hodnota robustního z-skóre v rámci daného parametru. Normalizace zajišťuje, že všechny parametry přispívají k výslednému skóre ve srovnatelném měřítku bez ohledu na jejich fyzikální jednotky nebo rozsah hodnot.

Výsledné anomaly score každého kusu je definováno jako aritmetický průměr normalizovaných z-skóre přes všechny sledované parametry:

$$\text{anomaly_score} = \frac{1}{n} \sum_{i=1}^n z_{norm,i}, \quad (3.3)$$

kde anomaly_score je výsledné skóre anomálnosti daného kusu, n je celkový počet sledovaných parametrů, $z_{norm,i}$ je normalizované z-skóre i -tého parametru nabývající hodnot v intervalu $(0, 1)$ a $\sum$ označuje součet přes všechny parametry $i = 1, 2, \dots, n$.

Je nutné zmínit **omezení vyplývající z povahy použité metody**: robustní z-skóre standardně předpokládá přibližně symetrické rozdělení hodnot parametru, což odpovídá normálnímu nebo obdobnému rozdělení; v praxi však část výrobních parametrů takový předpoklad nesplňuje – zejména parametry s přirozenou dolní nebo horní toleranční hranicí vykazují zpravidla logaritmicko-normální rozdělení, u něž je distribuce hodnot asymetrická. Za účelem zmírnění tohoto efektu bude, jak už bylo zmíněno v kapitole Předzpracování dat, na vybrané parametry **aplikována Yeo–Johnsonova transformace, která umožňuje transformovat i nenormálně rozdělená data** (včetně hodnot obsahujících nulu či záporné hodnoty) směrem k symetričtějšímu, přibližně normálnímu rozdělení (Yeo & Johnson, 2000). Tím se zlepšuje použitelnost robustního z-skóre pro tyto parametry a snižuje se riziko, že přirozená asymetrie dat bude mylně interpretována jako anomální chování.

Aplikace robustního z-skóre na takto upravené parametry je tedy **méně zatížena systematickým zkreslením**; přesto však mohou existovat případy, kdy ani transformace nevede k dostatečnému přiblížení k symetrickému rozdělení.

Výstupem této fáze bude rozšířený **dataset obsahující pro každý kus vypočtené anomaly score** a dílčí normalizovaná z-skóre pro vybrané parametry – stanovená granularita umožňuje nejen identifikaci anomálních kusů, ale také zpětnou analýzu toho, které parametry se nejvýrazněji podílely na jejich odchylce od standardního chování.

Na základě výše uvedeného lze formulovat konkrétní implikace pro další postup řešení. Navržený přístup **kombinuje řízenou redukci datového prostoru s následnou detekcí odchylek**, přičemž klíčovým prvkem je nahrazení čistě statistického výběru parametrů expertně řízenou selekcí. Tento krok reflektuje specifika výrobních dat, u nichž nízká variabilita neimplikuje nízkou informační hodnotu, a naopak potvrzuje nutnost zapojení doménové znalosti do analytického procesu.

3.4 Závěry



- Analýza bude probíhat nad **redukovaným, ale informačně relevantním podmnožinovým prostorem parametrů**, což umožní efektivnější a interpretovatelnější modelování.
- Zvolená metodika **robustního z-skóre, doplněná o normalizaci a případnou transformační úpravu dat**, poskytuje jednotné měřítko pro porovnání odchylek napříč heterogenními veličinami.
- Tím **vzniká kompaktní indikátor anomálosti**, který je dále využitelný jak pro identifikaci nestandardních kusů, tak pro jejich detailní rozklad na úroveň jednotlivých parametrů.
- Navržené řešení není pouze nástrojem pro jednorázovou analýzu, ale **představuje obecně aplikovatelný rámec pro průběžné vyhodnocování výrobních dat**. V následujících částech bude proto tato metodika prakticky implementována na reálných datech pilotní linky, přičemž důraz bude kladen nejen na identifikaci anomálií, ale také na posouzení přínosu navrženého přístupu z hlediska snížení objemu dat a potenciálních ekonomických dopadů.

4. Optimalizační řešení pro křivková data



Účelem je specifikovat optimalizační řešení pro křivková data, která jsou **specifickým typem výrobních záznamů a zachycují průběh fyzikálních veličin** během realizace jednotlivých technologických operací. Na rozdíl od skalárních parametrů **obsahují křivky detailní informaci o dynamice průběhu měření** – jsou nicméně generovány jen pro určité části výrobního procesu.

Navrhované **optimalizační řešení pro křivková data** se proto zaměřuje především na jejich efektivní reprezentaci a využití pro identifikaci nestandardních průběhů výrobního procesu. **Cílem je snížit objem ukládaných dat při zachování klíčové informační hodnoty** signálu a současně vytvořit mechanismus, který umožní automaticky rozpoznat anomální průběhy měření.

4.1 Předzpracování dat

Získaná křivková data byla zastoupena ve formátu JSON uloženém v textových souborech s příponou .txt, přičemž každý jednotlivý soubor reprezentoval průběh jedné operace na konkrétním výrobku v rámci výrobní linky. **Záznamy obsahují dvojice hodnot popisující průběh fyzikálních veličin**, které vznikají při rotačních operacích, typicky v pár tisících bodů s jemnou granularitou, díky které je následně možné křivku detailně opticky analyzovat. **Každý záznam se skládá ze dvou hlavních částí** – první jsou **metadata**, která obsahují základní popis měření, **následována samotnými daty** o měření, tedy hlavní měřenou částí.

Prvním krokem návrhu optimalizačního řešení je předzpracování dat, jehož **cílem je převést surové záznamy z výrobního procesu do konzistentní a analyticky využitelné podoby**. Vzhledem k tomu, že původní data jsou uložena v textových souborech ve formátu JSON a obsahují rozsáhlé množství metadat, je nutné provést jejich transformaci do strukturované formy.

Plánované kroky předzpracování zahrnují nejprve **kontrolu kvality vstupních souborů**, jejímž cílem je ověřit jejich úplnost, čitelnost a správnost struktury. Následně bude provedeno **parsování záznamů do formátu CSV**, které zajistí jednotnou tabulkovou reprezentaci dat pro jednotlivé měřené průběhy. Po převodu do formátu CSV bude následovat přepočít vybraných hodnot tak, aby byly křivky graficky lépe čitelné a vzájemně porovnatelné.

4.2 Komprese

Vzhledem k tomu, že křivková data představují z pohledu objemu i četnosti generování jednu z významnějších částí výrobních záznamů a **každá operace na výrobní lince vytváří průběhy měřených fyzikálních veličin, typicky o stovkách až tisících vzorků**, což při kontinuálním provozu vede k exponenciálnímu růstu objemu dat, představuje takový jev zásadní výzvu jak z hlediska požadavků na úložnou kapacitu a datové přenosy, tak i z pohledu následného zpracování a analýzy. Cílem této kapitoly je proto představit metody komprese, které umožňují reálně snížit datovou velikost křivek, aniž by došlo k významné ztrátě jejich informačního obsahu.

V reálném provozu by tato kompresní metoda nacházela uplatnění především ve fázi archivace dat. Současná praxe ve firmě spočívá v tom, že se všechny křivkové záznamy ukládají v plném, původním rozlišení, a to bez jakékoli redukce počtu vzorků. Dlouhodobě tak **narůstá objem historických dat**, která jsou sice důležitá pro zpětné dohledání výrobních událostí, ale pro naprostou většinu analytických nebo diagnostických úloh není nutné uchovávat je v maximálním detailu.

Po důkladné rešerši možností redukce bodů na křivce byl **zvolen algoritmus Douglas–Peucker**, který patří mezi nejrozšířenější metody pro zjednodušování trajektorií a datových křivek. Tento algoritmus byl původně vyvinut pro kartografické účely, avšak díky své univerzálnosti našel široké uplatnění také v průmyslových a signálních aplikacích. **Princip** tohoto algoritmu spočívá **v rekursivním rozdělávání trajektorie na lineární segmenty a postupném odstraňování bodů**, které leží uvnitř definované tolerance a nepřispívají k tvarové informaci signálu. Zachovány jsou pouze ty body, jejichž vynechání by vedlo k překročení stanoveného prahového odchýlení od původní křivky (Douglas & Peucker, 1973).

Základní myšlenka metody vychází z **geometrického hodnocení významnosti jednotlivých bodů**. Každý bod je posuzován podle své kolmice na spojnici mezi počátečním a koncovým bodem

segmentu. Pokud je vzdálenost bodu od této přímky menší než zvolený práh, je bod považován za redundantní a odstraní se. V opačném případě zůstává zachován a slouží jako nový uzlový bod pro další rekursivní dělení. Tímto způsobem postupně vzniká zjednodušená reprezentace signálu tvořená minimální množinou bodů, které stále popisují původní průběh s přijatelnou chybou aproximace (Douglas & Peucker, 1973).

Výsledkem je redukováný dataset, který si zachovává hlavní trend, extrémy a přechody křivky, zatímco drobné fluktuace jsou odstraněny. **Míra komprese** je přitom dána poměrem počtu bodů v komprimované a původní křivce. **Řídicí parametr algoritmu – tzv. prostorový práh**, značený ϵ – určuje kompromis mezi úrovní zjednodušení a věrností rekonstruovaného průběhu – volba optimální hodnoty tohoto parametru je zásadní pro dosažení rovnováhy mezi kvalitou aproximace a efektivitou komprese (Douglas & Peucker, 1973).

Vzhledem k tomu, že úspěšnost detekce anomálií je do značné míry závislá na zachování detailní informace o průběhu procesní křivky, je při návrhu experimentálního postupu nutné věnovat pozornost také pořadí jednotlivých kroků zpracování. **Kompresní metody**, přestože mohou výrazně snížit objem ukládaných dat, totiž zároveň **představují určitou formu aproximace původního signálu** a mohou tak **vést ke ztrátě jemných tvarových charakteristik**, které mohou být pro rozpoznání nestandardních průběhů klíčové. Z tohoto důvodu je vhodné detekční model nejprve navrhnout a ověřit na původních, nekomprimovaných datech, která zachovávají maximální informační obsah měření. Teprve po ověření funkčnosti navrženého přístupu lze následně analyzovat, jak se výsledky detekce mění při aplikaci kompresních metod. Posouzení vlivu komprese na kvalitu detekce anomálií tak představuje potenciální směr dalšího výzkumu, který může přispět k nalezení rovnováhy mezi úsporou datového objemu a zachováním diagnostické hodnoty procesních křivek.

4.3 Analýza významnosti

V této fázi návrhu optimalizačního řešení bude pozornost zaměřena na určení **informační významnosti jednotlivých úseků procesní křivky**, která popisuje vztah mezi momentem a úhlem. Cílem je určit, které části průběhu mají zásadní vliv na charakteristiku procesu a mohou sloužit k rozpoznání odchylek, a které naopak obsahují převážně opakující se nebo málo přínosné informace.

Navrhovaný přístup vychází z principů používaných při **analýze časových řad a rozpoznávání tvarových vzorů**, přestože v tomto případě nejsou data závislá na čase, ale na úhlu, což znamená, že místo sledování vývoje veličiny v čase je sledována trajektorie procesu – tedy průběh momentu vzhledem k měnícímu se úhlu (Olszewski, 2001).

Pro určení významnosti jednotlivých úseků se uvažuje **rozdělení křivky na menší části**, z nichž každá představuje určitý lokální tvarový úsek. Úseky lze popsat prostřednictvím několika základních charakteristik, například rozsahem hodnot, proměnlivostí, sklonem nebo plochou pod křivkou; tyto kvantitativní ukazatele dále umožní srovnávat jednotlivé úseky mezi sebou a posoudit, jaký mají přínos k celkovému průběhu (Olszewski, 2001).

Každý z těchto úseků bude vyjádřen pomocí souboru číselných rysů, jež popíší jeho hlavní vlastnosti. Hodnoty budou následně převedeny do společného měřítka, aby bylo možné porovnávat jejich relativní význam napříč celou křivkou. Tím vznikne jednotný rámec pro určení, které části mají největší diagnostickou hodnotu (Olszewski, 2001).

Součástí návrhu je také porovnání s expertním posouzením, které ověří, **zda výsledky vypočtené podle navržené metody odpovídají reálné technologické důležitosti** jednotlivých fází procesu. Při rozcházejících se názorech umožní tento krok případně upravit váhy jednotlivých ukazatelů tak, aby výsledná analýza lépe vystihovala skutečné chování výrobního systému (Olszewski, 2001).

4.4 Detekce anomálií

Cílem této části je **formulovat postup detekce anomálií v procesních křivkách**, který by umožnil včasnou identifikaci nestandardních průběhů výrobních operací a přispěl tak k prevenci vzniku vadných kusů i následných reklamačních řízení.

Dosud používané metody průmyslového monitoringu vycházejí převážně z pevně stanovených prahových hodnot nebo z okénkové analýzy, která sleduje vybrané charakteristiky signálu v časově či úhlově definovaných úsecích – přístupy jsou sice výpočetně nenáročné a dobře interpretovatelné, avšak jejich schopnost odhalovat složitější či nelineární odchylky je omezená a do značné míry závislá na zkušenostech obsluhy. S rostoucí komplexitou výrobních procesů a objemem dostupných dat se proto

ukazuje **potřeba využívat pokročilejší metody**, které umožňují zachytit jemné vzory a vztahy v datech i bez nutnosti jejich předchozí explicitní specifikace (Tan et al., 2005).

V rámci této kapitoly je pozornost zaměřena především na **přístupy založené na strojovém a hlubokém učení, zejména na jednotřídní modely**, které se učí pouze z běžných průběhů a následně vyhodnocují jejich odchylky od standardního chování. Přístup se ukázal vhodným zejména díky existenci omezeného množství anomálních záznamů a faktu, že jejich systematické získávání by bylo časově, a tedy i ekonomicky, neefektivní.

Následující text nejprve **shrnuje tradiční metody monitoringu** a jejich limity, poté se věnuje moderním přístupům k detekci anomálií v procesních křivkách a konečně představuje návrh experimentálního postupu, který bude aplikován na dostupná výrobní data.

V průmyslové praxi se historicky **využívají konvenční metody sledování procesních křivek**, které jsou založeny na předem definovaných parametrech a explicitně stanovených pravidlech. Mezi nejčastěji používané patří okénková analýza, při níž je křivka rozdělena do několika časových nebo jiných segmentů z nichž v každém se vyhodnocují vybrané charakteristiky – průměrná hodnota, maximum nebo třeba rychlost změny signálu (Tan et al., 2005).

Dalším rozšířeným postupem je **kontrola vůči tolerančním mezím**, kdy se průběh sledované veličiny porovnává s horní a dolní limitní hodnotou, jejichž překročení signalizuje možnou odchylku procesu. Specifickou variantu představuje porovnávání s referenční „obálkovou“ křivkou, kdy je měřený průběh vyhodnocován vzhledem k etalonové křivce obklopené tolerančním pásmem určujícím přípustnou variabilitu signálu (Tan et al., 2005).

Ačkoli jsou tyto metody v průmyslu dlouhodobě využívány, jejich účinnost je podmíněna pečlivou kalibrací a vysokou mírou a konzistencí expertních znalostí personálu. V praxi navíc vykazují omezenou schopnost detekovat složitější, nelineární či subtilní odchylky od standardního průběhu, jež mohou být pro kvalitu výrobku kritické, přičemž neporušují explicitně stanovené limity. Tato **omezení jsou obzvláště patrná v prostředí moderní výroby, kde procesní data vykazují vysokou variabilitu**, jsou ovlivněna více faktory současně a mohou obsahovat odchylky, jejichž vzory nelze předem plně definovat (Tan et al., 2005).

S postupující digitalizací průmyslové výroby a současným důrazem na efektivní snižování nákladů spojených s ukládáním a zpracováním dat dochází **k významnému rozšíření možností aplikace pokročilých analytických metod nad procesními křivkami**, a to zejména s využitím nástrojů strojového a hlubokého učení. V kontextu montážních a spojovacích operací, mezi které patří právě i šroubování nebo třeba lisování, představují procesní křivky – typicky záznam průběhu momentu utahování v čase nebo závislost síly na dráze – zdroj informací pro zajištění kvality výroby a včasnou detekci procesních odchylek či poruch (Meiners et al., 2020).

V odborné literatuře lze nalézt široké **spektrum metod pro detekci anomálií** v procesních křivkách, přičemž volba konkrétního přístupu závisí na dostupnosti dat, povaze výrobního procesu a požadované citlivosti na odchylky. Z hlediska konstrukce modelu se tyto metody dělí na jednotřídní a vícetřídní.

Jednotřídní metody jsou v průmyslové praxi často preferovány, neboť defektní průběhy jsou obvykle vzácné a jejich systematické sbírání je nákladné či technologicky obtížné. Modely tohoto typu jsou trénovány výhradně na záznamech odpovídajících normálnímu chodu procesu a jejich úkolem je vytvořit reprezentaci „standardního“ chování. Odchylky od tohoto modelu jsou následně klasifikovány jako anomálie. Mezi běžně využívané **metody patří One-Class Support Vector Machines, autoenkodéry a Isolation Forest**. **One-Class Support Vector Machines** konstruuje rozhodovací hranici kolem dat normální třídy, autoenkodéry využívají neuronové sítě k rekonstrukci vstupních dat a anomálie identifikují na základě velikosti rekonstrukční chyby, zatímco Isolation Forest používá náhodné rozdělení dat k rychlé izolaci odlehklých průběhů (Meiners et al., 2020).

V situacích, kdy je k dispozici dostatek jak normálních, tak defektních záznamů, je možné využít vícetřídní klasifikaci. Ta má za cíl přímo **oddělit jednotlivé třídy průběhů a často se realizuje pomocí Support Vector Machines, Random Forests, Gradient Boosted Trees nebo neuronových sítí**, například vícevrstvých perceptronů a konvolučních neuronových sítí. Konvoluční neuronové sítě umožňují extrahovat tvarové rysy přímo z křivek nebo z jejich transformací, například do časově-frekvenčních reprezentací (spektrogramů) (Meiners et al., 2020).

Vedle čistě modelových řešení se uplatňují i hybridní přístupy, které **kombinují automatické učení s ruční extrakcí příznaků** – mezi typické patří maximální dosažený moment, úhel při dosažení určité

síly či celková doba trvání operace. Tento přístup je výhodný zejména u menších datasetů, protože poskytuje interpretovatelné výstupy a může zlepšit výkon modelu (Meiners et al., 2020).

Dle odborné literatury se ukazuje, že **jednotřídní modely mají ve výrobním prostředí zásadní praktickou výhodu oproti víceřídním přístupům**, zejména kvůli nedostatku kvalitně anotovaných dat o defektech. Jejich předností je schopnost detekovat i zcela nové typy odchylek, které nebyly v tréninkových datech obsaženy, a lepší adaptace na proměnlivé výrobní podmínky. Daný **vývoj směřuje k metodám, které kombinují vysokou citlivost detekce s možností implementace v reálném čase**, při optimalizaci výpočetní náročnosti, latence a kompatibility s existujícími systémy.

Vzhledem k povaze dostupných dat se jako výchozí strategie jeví jednotřídní učení nad „zdravými“ průběhy. Dostupných 502 OK záznamů tvoří dostatečný základ k naučení reprezentace standardního chování procesu, aniž by bylo nutné sbírat rozsáhlé množství defektních vzorků. Sada 60 evaluačních záznamů, která zahrnuje **směs OK a anomálií, bude sloužit k nezávislému nastavení prahové hodnoty, ladění hyperparametrů a posouzení generalizace** bez úniku informací z anomálií do tréninku.

V první fázi se předpokládá sjednocení dat z JSON formátu do časových křivek pevné délky a jejich normalizace na srovnatelnou škálu, aby **byly odstraněny triviální rozdíly způsobené měřítkem a vzorkováním**. Následně by byl natrénován jednotřídní model, který se naučí rekonstruovat či ohraničit prostor běžných průběhů. Přirozeným kandidátem je autoenkodér nad 1D křivkami (pro zachycení nelineárních vzorů bez nutnosti ruční extrakce příznaků), zatímco metoda Isolation Forest se vzhledem k povaze vstupů nezdá být velmi vhodnou, a tak není její využití v plánu.

Využití metod vyžadujících **vyvážené množství tréninkových dat obou tříd se za daných podmínek nepředpokládá** – víceřídní klasifikační modely by totiž v takovém případě mohly být náchylné k přeučení, a navíc by postrádaly schopnost detekovat nové typy odchylek, které se nenachází v současném tréninkovém datasetu.

Rozhodovací práh pro rozlišení mezi normálním a anomálním průběhem bude určen empiricky z trénovacích dat – bude po expertní validaci definován jako 99. percentil rekonstrukční chyby dosažené na trénovací sadě obsahující relativně bezvadné křivky. Přístup reflektuje přirozenou variabilitu i v rámci bezvadných průběhů, kdy část křivek může vykazovat vyšší rekonstrukční chybu, aniž by se jednalo o skutečnou anomálii.

4.5 Závěry



- **Hlavní problém** nespočívá pouze v jejich analýze, ale už **v samotném objemu a formě**, ve které jsou uchovávána.
- Navržený postup proto **odděluje práci s daty na dvě roviny**:
 - na jedné straně řeší jejich efektivní reprezentaci a možnou redukci,
 - na straně druhé jejich využití pro detekci odchylek.
- Nezanedbatelné je přitom **zachování tvarové informace křivky**, která nese podstatnou diagnostickou hodnotu, a zároveň omezení redundantních detailů, jež nejsou pro další analýzu nezbytné.
- **Detekce anomálií** je následně uvažována jako **problém rozpoznání odchylky** od běžného průběhu, přičemž zvolený přístup reflektuje omezenou dostupnost defektních dat a opírá se primárně o modelování standardního chování procesu.

5. Aplikace navržených metod pro skalární data



Účelem je charakterizovat aplikace pro zpracování skalárních dat, a to na základě předem navržených metod.

5.1 Předzpracování dat

Zpracování zaslaných dat vycházelo z charakteru dostupných výrobních záznamů – obdržená složka obsahovala data ze všech dostupných výrobních stanic pro danou výrobní linku, přičemž každá ze stanic byla reprezentována jedním CSV souborem z odpovídajícího časového období; časová značka se pro stejný identifikátor zpravidla o pár vteřin lišila kvůli skutečnosti, že každá stanice má vlastní časový senzor. Jednotlivé **soubory zpravidla v prvních sloupcích obsahovaly metadata**, následně pak **měřené hodnoty parametrů** a případně i jejich **dolní a horní toleranční meze**.

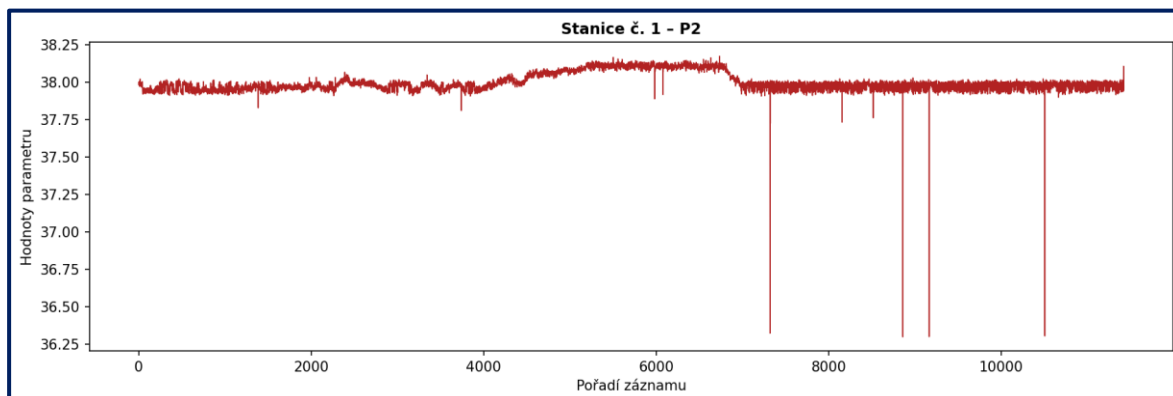
Před započítím analýzy byla provedena série kroků pro zajištění kvality a konzistence dat – nejprve byly **odstraněny záznamy obsahující nadměrné množství chybějících hodnot**. Vzhledem k cíli analýzy byly dále odstraněny všechny záznamy, které byly označeny jinak než „good“, tedy OK. Při další exploraci se ukázalo, že v datech dle očekávání **dochází k opakovanému výskytu záznamů** vztahujících se ke stejnému identifikátoru; takové případy byly ošetřeny selektivní redukcí záznamů, přičemž za reprezentativní byl považován ve většině případů vzorek odpovídající nejnovějšímu časovému okamžiku měření, a tak bylo zajištěno, že každý výrobní kus byl v datové sadě zastoupen právě jednou hodnotou, aniž by došlo ke ztrátě relevantní informace.

5.2 Selektce parametrů a jejich popis

Výběr parametrů pro následnou analýzu byl realizován **skrze expertní selekci** – proces probíhal formou konzultací s procesními specialisty disponujícími detailní znalostí fyzikální podstaty měřených veličin i technologických omezení jednotlivých výrobních stanic. Na základě jejich posouzení bylo vybráno **šest výrobních stanic s konkrétními parametry**, přičemž dvě ze stanic byly označeny jako nejdůležitější z hlediska identifikace nestandardního chování výrobního procesu.

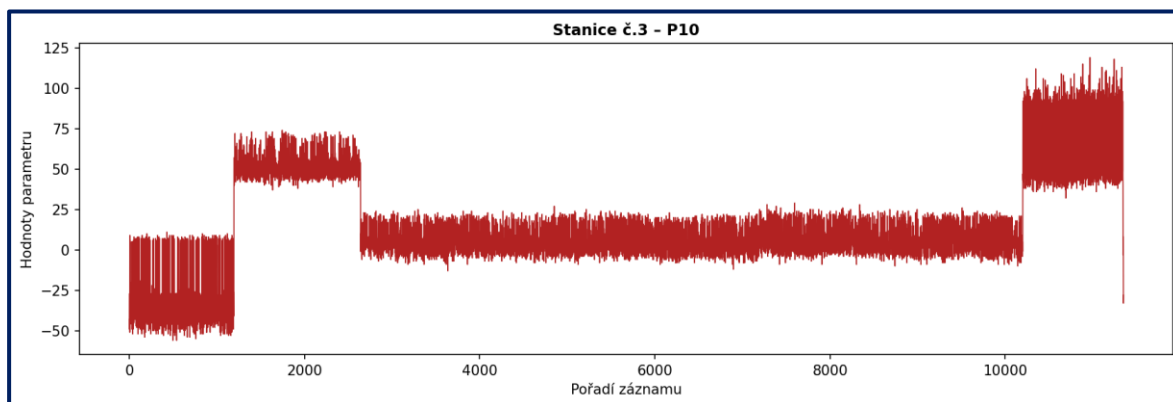
Součástí selekčního procesu byla také **vizuální explorace dat**. Pro jednotlivé parametry byly analyzovány grafy časových řad a histogramy rozložení hodnot, které umožnily lépe porozumět charakteru dat a identifikovat případné extrémní hodnoty. V některých případech byly pozorovány **odlehle hodnoty technologického původu**, tedy hodnoty, které se sice statisticky jeví jako extrémní, avšak z hlediska fyzikálního principu měření odpovídají běžnému chování výrobního procesu. Takové parametry nebyly automaticky vyloučeny z analýzy, ale byly dále posuzovány individuálně s ohledem na jejich technologický význam.

Celkem bylo expertem do analýzy vybráno 112 parametrů různých rozdělení. Po vizuální exploraci byly vyřazeny tyto parametry:



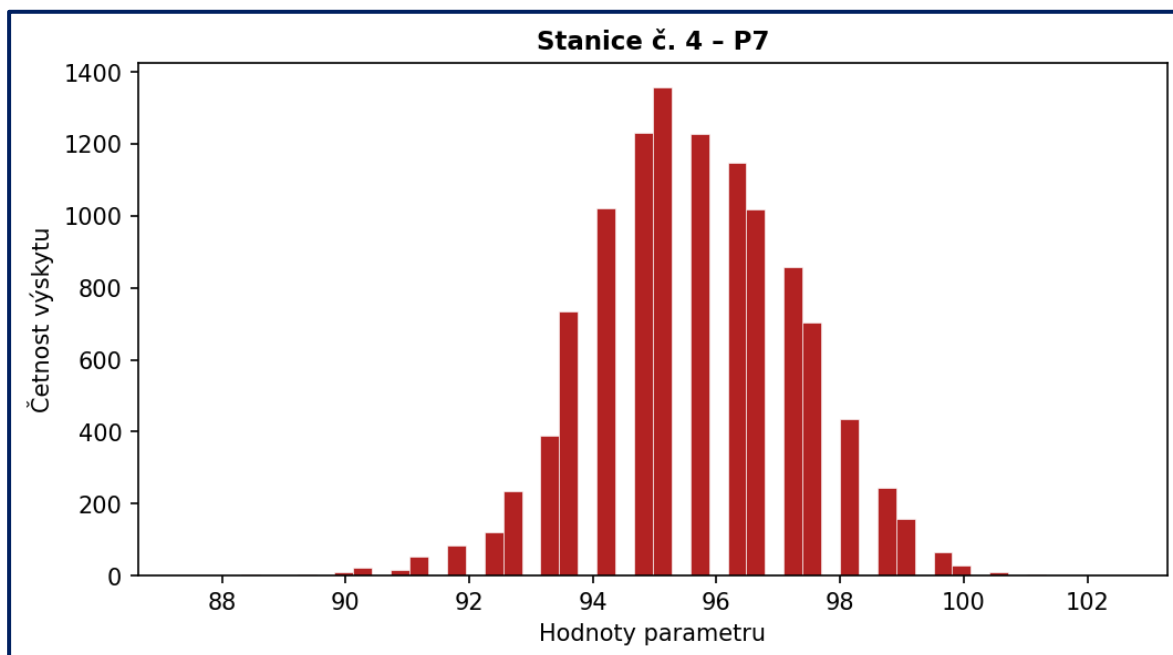
Obrázek 5.1: Časová řada druhého parametru na stanici č. 1 (zdroj: autorka, 2026)

Časová řada druhého parametru na stanici č. 1 vykazuje **nestabilitu hodnot** v průběhu sledovaného období. Přestože se většina měření pohybuje v relativně úzkém rozsahu, v datech se objevuje postupný nárůst a následný pokles hodnot, které nejsou doprovázeny žádnou systematickou změnou trendu ani opakujícím se vzorem – výkyvy se vyskytují poněkud nepravidelně a výsledek skóre na tomto parametru by byl zavádějící.



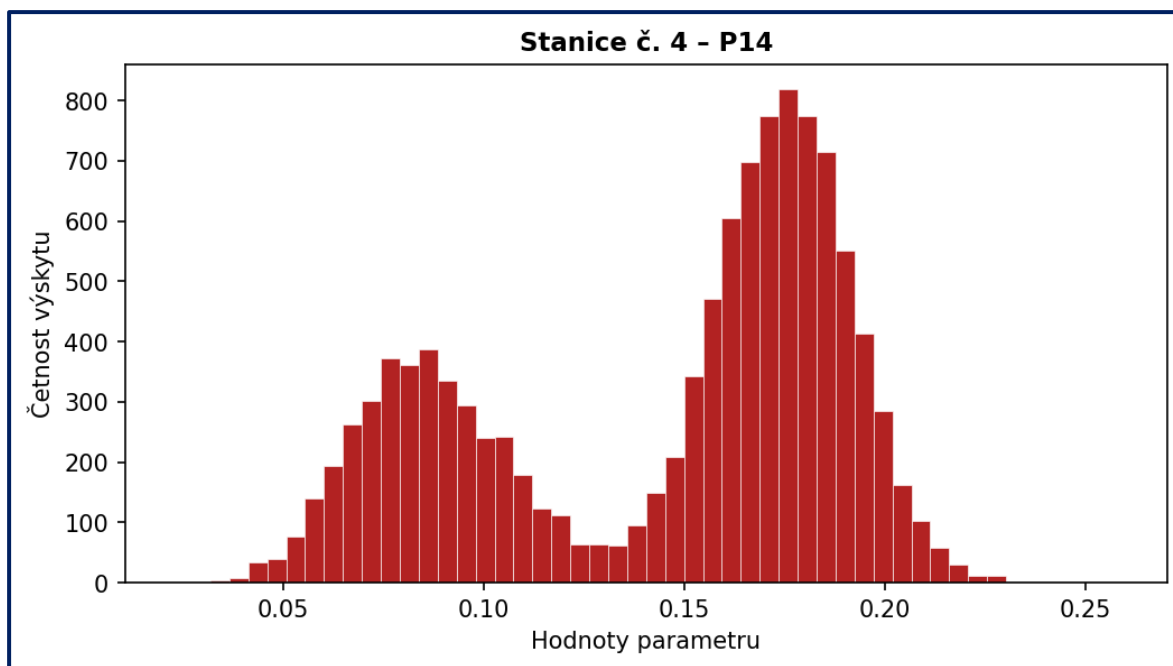
Obrázek 5.2: Časová řada desátého parametru na stanici č. 3 (zdroj: autorka, 2026)

Časová řada desátého parametru na stanici č. 3 vykazuje výrazné změny úrovně hodnot v průběhu sledovaného období. Data se postupně pohybují ve třech relativně stabilních pásmech, mezi nimiž dochází k náhlým skokovým přechodům. **Změny nejsou plynulé, ale mají charakter náhlých posunů** celé úrovně měření. Za takovýchto podmínek nelze aplikovat robustní z-skóre a očekávat spolehlivý výsledek.



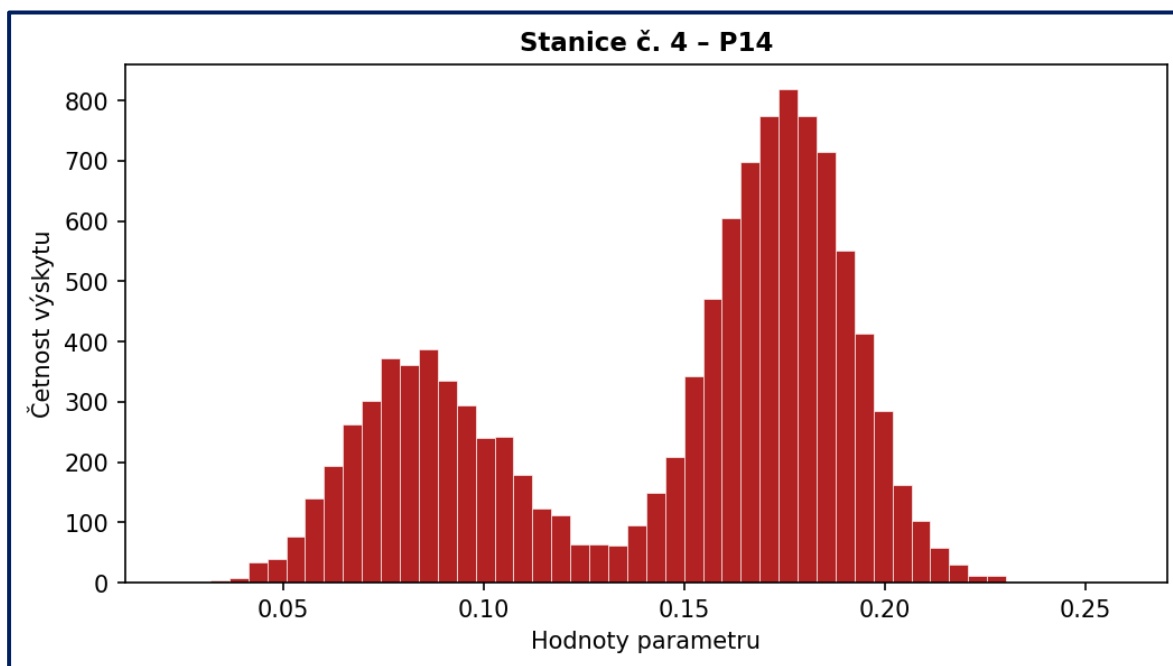
Obrázek 5.3: Histogram četností sedmého parametru na stanici č. 4 (zdroj: autorka, 2026)

Histogram sedmého parametru na stanici č. 4 ukazuje diskrétní charakter rozdělení hodnot, který vyplývá z principu měřicího zařízení umožňujícího zaznamenat pouze omezené spektrum hodnot. **Obdobné rozdělení bylo identifikováno také u dalších sedmi parametrů** na některých výrobních stanicích. Tyto případy proto nejsou z důvodu přehlednosti dále samostatně ilustrovány.



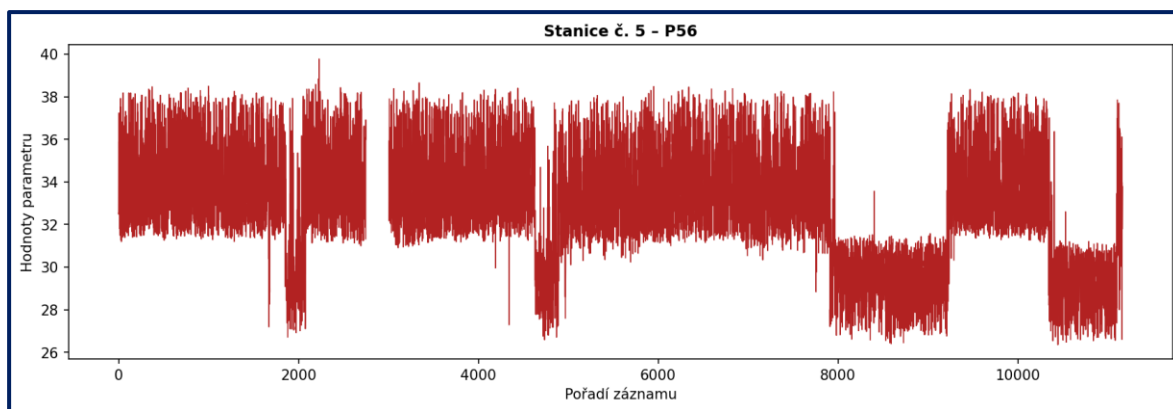
Obrázek 5.4: Histogram četností čtrnáctého parametru na stanici č. 4 (zdroj: autorka, 2026)

Histogram čtrnáctého parametru na stanici č. 4 vykazuje výrazně bimodální rozdělení hodnot, kdy se naměřené hodnoty koncentrují do dvou oddělených intervalů, kvůli přítomnosti dvou odlišných provozních stavů. Obdobné bimodální rozdělení bylo pozorováno také u dalších dvou parametrů na jiných výrobních stanicích. Z důvodu přehlednosti však tyto případy nejsou dále samostatně graficky prezentovány.



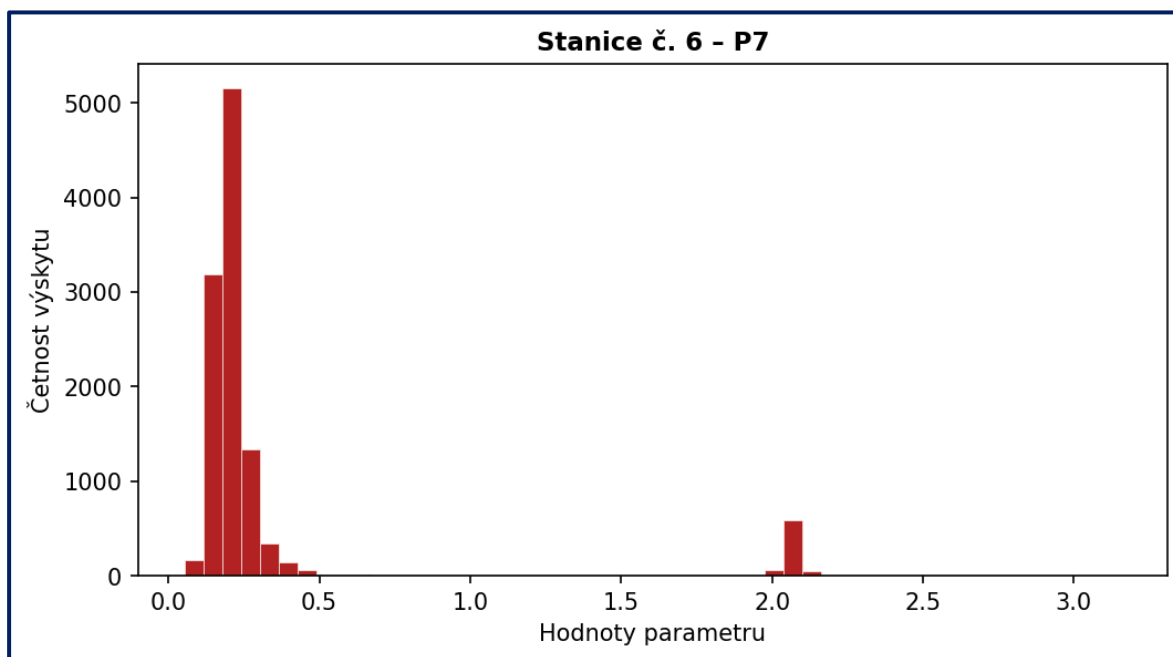
Obrázek 5.5: Histogram četností čtyřicátého osmého parametru na stanici č. 5 (zdroj: autorka, 2026)

Histogram čtyřicátého osmého parametru na stanici č. 5 vykazuje trimodální rozdělení hodnot, kdy se naměřené hodnoty soustřeďují do několika oddělených intervalů, a je tedy pro další analýzu naprosto nevhodný. Obdobný typ rozdělení byl pozorován také u dalšího parametru na stejné výrobní stanici; z důvodu přehlednosti už není znovu zobrazován.



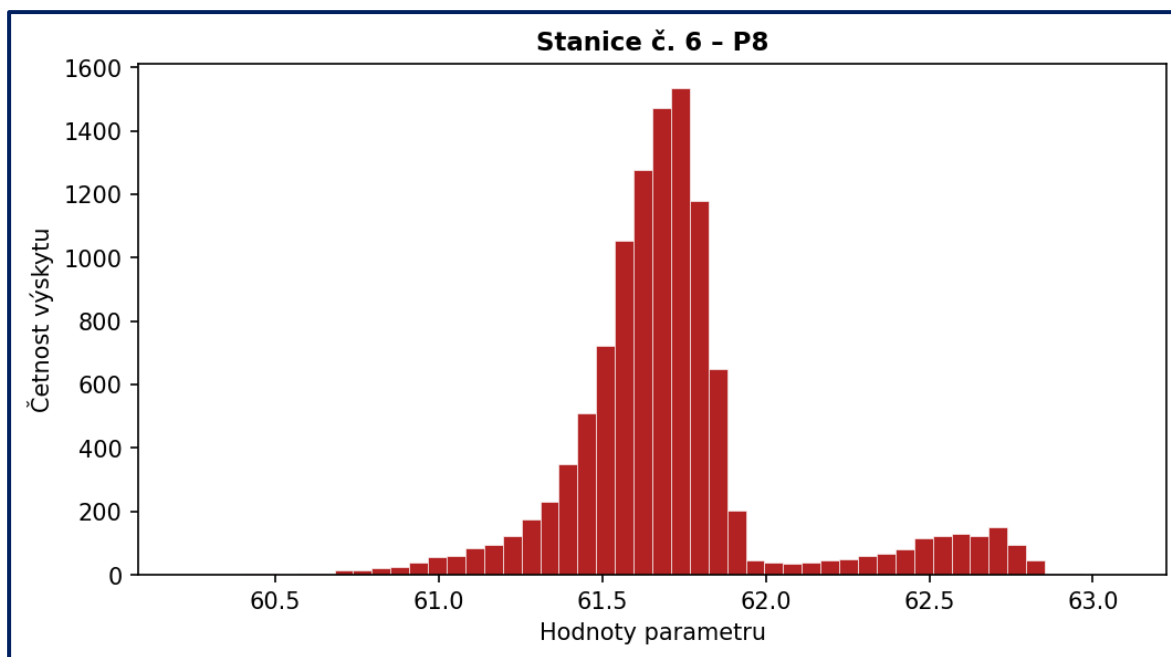
Obrázek 5.6: Časová řada padesátého šestého parametru na stanici č. 5 (zdroj: autorka, 2026)

Časová řada padesátého šestého parametru na stanici č. 5 vykazuje výrazné změny úrovně hodnot v průběhu sledovaného období. Naměřené hodnoty se střídavě pohybují v několika relativně stabilních pásmech, mezi nimiž dochází k náhlým skokovým přechodům. Vzhledem k tomu, že parametr nevykazuje stabilní statistické vlastnosti v čase a jeho chování je výrazně ovlivněno těmito skokovými změnami, byl z další analýzy vyřazen.



Obrázek 5.7: Histogram četností sedmého parametru na stanici č. 6 (zdroj: autorka, 2026)

Histogram sedmého parametru na stanici č. 6 vykazuje výrazně asymetrické rozdělení hodnot – převážná většina měření se koncentruje v úzkém intervalu nízkých hodnot, zatímco menší část pozorování se nachází ve výrazně vzdálené oblasti vyšších hodnot, což vede k silné nerovnováze mezi hlavní koncentrací dat a izolovanými hodnotami, které tvoří samostatný shluk. Toto rozdělení komplikuje interpretaci odlehklých hodnot a může vést k jejich nesprávné identifikaci při použití statistických metod detekce anomálií. Z tohoto důvodu byl parametr z další analýzy vyřazen.



Obrázek 5.8: Histogram četností osmého parametru na stanici č. 6 (zdroj: autorka, 2026)

Histogram osmého parametru na stanici č. 6 vykazuje specifické bimodální rozdělení hodnot s dominantní koncentrací měření v jednom hlavním intervalu. Vedle tohoto hlavního shluku se v datech objevuje menší skupina hodnot ve výrazně vyšší oblasti, která vytváří sekundární koncentraci dat. Vzhledem k tomu, že rozdělení hodnot není homogenní a obsahuje oddělené skupiny pozorování, může být identifikace anomálií u tohoto parametru statisticky zkreslená. Parametr byl proto z další analýzy vyřazen.

Po vyřazení parametrů s nestandardním rozdělením bylo do analýzy **zařazeno 94 parametrů s normálním nebo logaritmicke-normálním rozdělením**.

5.3 Detekce anomálií na skalárech

Detekce anomálií probíhala **samostatně pro každou výrobní stanici** a pro předem definovaný seznam parametrů vybraných na základě expertní selekce a následné korekce. Po načtení dat byly jednotlivé parametry převedeny na numerický formát.

U každého analyzovaného parametru byla nejprve **posouzena šikmost jeho rozdělení**. V případech, kdy **rozdělení vykazovalo výraznou asymetrii**, byla testována Yeo–Johnsonova transformace a byla aplikována pouze tehdy, pokud po jejím použití došlo ke snížení absolutní hodnoty šikmosti oproti původnímu rozdělení. Pokud transformace rozdělení nezlepšila, byl parametr ponechán v původní podobě.

Následně byla pro každý parametr **vypočtena odchylka jednotlivých hodnot** od typického chování pomocí robustního z-skóre.

Robustní z-skóre bylo pro jednotlivé parametry stanoveno podle již zmíněného vztahu:

$$z = \frac{(x - \text{medián}(x))}{MAD(x)}, \quad (5.4)$$

Vzhledem k tomu, že obecným cílem této metody nebylo rozlišovat směr odchylky, ale její velikost, byla použita absolutní hodnota robustního z-skóre. Vysoké kladné i záporné odchylky tedy byly interpretovány jako stejně významné z hlediska možného anomálního chování.

$$z^{norm} = \frac{|z|}{\max(|z|)}, \quad (5.5)$$

Pro zajištění vzájemné srovnatelnosti parametrů byla **absolutní hodnota robustního z-skóre normalizována** vydělením maximální hodnotou dosaženou v rámci daného parametru, čímž byl každý parametr převeden na srovnatelnou škálu v intervalu od 0 do 1, kde hodnota 0 odpovídá minimální odchylce, tedy střední hodnotě, a hodnota blíží se 1 nejvyšší pozorované odchylce v rámci daného parametru. Standardně se totiž parametry od sebe, zvláště v kontextu různých výrobních stanic, liší jak fyzikální jednotkou, tak i rozsahem a rozptylem hodnot.

Po normalizaci byla pro každý výrobní kus **stanovena celková hodnota anomaly score** na výrobní stanici jako aritmetický průměr normalizovaných absolutních robustních z-skórů napříč všemi zahrnutými parametry skrze následující, již zmíněný vzorec:

$$anomaly_score = \frac{1}{n} \sum_{i=1}^n z_{norm,i}. \quad (5.6)$$

Takto definované skóre **reprezentuje souhrnnou míru odchylky konkrétního výrobního kusu** od běžného chování procesu. Čím vyšší je jeho hodnota, tím výrazněji se daný kus odlišuje od typického průběhu sledovaných parametrů. Výhodou tohoto přístupu je skutečnost, že anomálie není posuzována pouze na základě jednoho izolovaného parametru, ale jako kombinovaný efekt více veličin současně.

Výsledkem detekce byl pro každou analyzovanou stanici **výstupní datový soubor obsahující původní záznamy doplněné o dílčí normalizovaná z-skóre** jednotlivých parametrů a o výsledné celkové anomaly score. Tento výstup následně sloužil jako podklad pro identifikaci podezřelých výrobních kusů a pro další interpretaci z hlediska technologického významu odhalených odchylek.

Tabulka 5.1: Ukázka výsledného anomaly score na stanici č. 2 (zdroj: autorka, 2026)

P1	P2	anomaly_score
0,03287	0,00527	0,01907
0,01233	0,00763	0,00998
0,05245	0,01224	0,03235
0,01058	0,00047	0,00552

Po výpočtu anomaly score na úrovni jednotlivých výrobních stanic byly výsledky dále propojeny do společného přehledu, aby bylo možné posoudit míru odchylky výrobního kusu v širším kontextu celého sledovaného úseku výroby. **Pro každou stanici byl využit výstupní soubor obsahující identifikátor výrobního kusu a odpovídající hodnotu anomaly score.**

Vzhledem k tomu, že v průběhu výrobního procesu dochází ke změně používaného identifikátoru, bylo nejprve nutné zajistit návaznost mezi starším a novějším formátem identifikace výrobků. Čistě k tomuto účelu byla využita data ze stanice č. 7 ve které byly současně dostupné oba typy identifikátorů a posloužila tak jako přemostující **článek propojující výsledky z předchozích a následujících výrobních stanic do jednoho společného souboru.**

Následně byly k jednotlivým **výrobním kusům přiřazeny hodnoty anomaly score** ze zvolených stanic. Z výsledného přehledu byly odstraněny záznamy, u nichž nebylo možné jednoznačně dohledat oba identifikátory, případně obsahovaly neplatné nebo chybějící hodnoty identifikace. Současně byly

ponechány pouze ty kusy, pro které byla dostupná alespoň hodnota anomaly score ze stanice č. 1, kvůli případné zavádějící neúplnosti výsledku.

Za účelem celkového posouzení míry odchylky výrobního kusu napříč více stanicemi bylo následně stanoveno agregované skóre jako **vážený průměr dílčích anomaly score**. S ohledem na počáteční expertní analýzu byly jednotlivým stanicím přiřazeny váhy reflektující jejich relativní význam při výpočtu **agregovaného anomaly score** – jejich výše byla opět stanovena na základě expertního vyjádření. Stanice č. 4 a stanice č. 5 tak dostaly zvýšenou váhu deset, zatímco stanice č. 1, 2, 3 a 6 dostaly přiřazenou standardní váhu 1. Stanice č. 7 dostala jako čistě přemostující váhu 0 a její parametry nebyly vůbec ve výpočtu zohledněny. **Výsledné agregované skóre tak lze vyjádřit vztahem**

$$\text{avg}_{\text{score}_j} = \frac{\sum_{s=1}^m w_s \cdot \text{score}_{j,s}}{\sum_{s=1}^m w_s}, \quad (5.7)$$

kde $\text{score}_{j,s}$ představuje hodnotu anomaly score j -tého výrobního kusu na s -té stanici, w_s odpovídá váze dané stanice a m je počet zahrnutých stanic, pro něž byla u daného kusu dostupná hodnota skóre. Do jmenovatele vstupují pouze váhy těch stanic, u nichž byla hodnota anomaly score skutečně k dispozici.

Takto definované agregované skóre **představuje souhrnnou míru odchylky výrobního kusu v rámci více navazujících částí výrobního procesu** – hlavní výhodou je, že umožňuje zohlednit nejen izolované odchylky na jednotlivých stanicích, ale i jejich případnou kumulaci v průběhu výroby. Současně lze prostřednictvím vah zdůraznit ty stanice, které jsou z technologického nebo analytického hlediska považovány za významnější.

Tabulka 5.2: Ukázka výsledku výpočtu anomaly score – seřazené od nejvyššího, bez označení identifikátory (zdroj: autorka, 2026)

Stanice č. 1	Stanice č. 2	Stanice č. 3	Stanice č. 4	Stanice č. 5	Stanice č. 6	Anomaly score
0,126	0,051	0,368	0,551	N/A	N/A	0,466
0,233	0,037	0,224	0,510	0,489	0,269	0,448
0,142	0,061	0,309	0,428	0,529	0,285	0,432
0,243	0,071	0,245	0,490	N/A	N/A	0,420
0,150	0,029	0,094	0,467	0,484	0,229	0,417
0,189	0,046	0,336	0,448	0,449	0,416	0,415
0,169	0,045	0,282	0,460	0,428	0,335	0,404
0,155	0,080	0,489	0,433	0,419	0,384	0,401

Výsledkem této fáze zpracování je **přehledový datový soubor** ve formátu Excel, v němž byl pro každý výrobní kus uveden **původní i navazující identifikátor, dílčí anomaly score z jednotlivých stanic a výsledná agregovaná hodnota skóre** – jak lze ostatně vidět na výše zobrazené tabulce (viz Tabulka 5.2). Tabulka zobrazuje prvních osm generovaných záznamů, které jsou seřazeny od nejvyššího k nejnižšímu dle vypočteného skóre. Obsahuje postupný vývoj hodnot anomaly score skrze stanice, kterými chronologicky prochází s výslednou průměrnou hodnotou, tedy zobrazuje jakýsi časový vývoj anomaly score při průchodu linkou. Ty záznamy, které mají na některých pozdějších stanicích (a hlavně na šesté) stanici hodnotu skóre označenou jako N/A jsou takové, které byly v některé fázi procesu označeny standardními metodami jako NOK. Na základě expertního posouzení byly takové kusy ponechány v celkovém přehledu, protože na nich lze zkoumat vývoj jejich anomaly score před vypadnutím. Naopak ty záznamy, které mají skóre vypsáno na všech stanicích, prošly standardně celým procesem výroby a byly odeslány.

Vzhledem k tomu, že nebyla k dispozici jednoznačná identifikace konkrétních reklamovaných kusů, nebylo možné provést standardní kvantitativní vyhodnocení přesnosti navržené metody pomocí klasických metrik – **hodnocení přínosu detekce anomálií bylo proto založeno primárně na expertním posouzení výsledků**. Na pilotní lince č. 1 se ročně vyskytuje pouze velmi malé množství reklamací, cílem navrženého přístupu tedy není jednoznačná klasifikace kusů na OK a NOK, ale identifikace těch, které se svým chováním odlišují od běžného průběhu výroby.

Z tohoto důvodu **nebyl v rámci metody stanoven pevný rozhodovací práh anomaly score**; výsledná hodnota skóre slouží výhradně jako nástroj pro seřazení výrobních kusů podle míry jejich odchylky od typického chování. Přístup tak odpovídá spíše principu prioritizace než klasifikace – místo snahy o jednoznačné rozhodnutí je cílem usnadnit výběr omezené množiny potenciálně problematických případů pro následnou kontrolu.

5.4 Závěry



- Expertní posouzení výsledného přehledu ukázalo, že **kusy s nejvyšší hodnotou anomaly score skutečně vykazují určité nestandardní charakteristiky** v průběhu sledovaných parametrů.
- Ačkoliv se nejedná ve všech případech o jednoznačně vadné výrobky, **vybrané kusy** představují z hlediska další analýzy **relevantní kandidáty**, jimž se vyplatí věnovat zvýšenou pozornost.
- Metoda tak umožňuje **zaměřit kontrolu na velmi malou část produkce** způsobem, který je výrazně efektivnější než náhodný výběr, neboť upřednostňuje kusy s nejvyšší mírou odchylky od běžného chování.
- Práce se skalárními výrobními daty stojí především **na jejich správném „očišťení“ a pochopení kontextu** – teprve po odstranění nekonzistencí, duplicit a problematických parametrů dává smysl přistoupit k samotné analýze.
- Výsledky současně potvrzují, že **chování jednotlivých veličin nelze posuzovat izolovaně ani čistě statisticky**; bez znalosti technologie by část parametrů byla vyhodnocena zavádějícím způsobem.
- Použitý způsob výpočtu odchylky **pak umožňuje převést různorodé veličiny na společnou škálu a porovnávat je mezi sebou**, přičemž největší přínos se ukazuje ve chvíli, kdy jsou dílčí odchylky spojeny napříč více stanicemi – v takovém případě již nejde pouze o identifikaci extrémů v jednotlivých měřeních, ale o zachycení kusů, u nichž se odchylky kumulují v průběhu výroby.

6. Aplikace navržených metod pro křivková data



Účelem je charakterizovat aplikace pro zpracování křivkových dat, a to na základě předem navržených metod..

6.1 Předzpracování a čištění dat

Každý z poskytnutých záznamů ve formátu JSON obsahoval **detailní průběh dotažení uložený v poli tightening steps**. Jednotlivé kroky měření zahrnují objekt graph, který obsahuje průběhy veličin – **angle values, torque values, gradient values a time values**. **Kromě těchto hlavních polí se v souborech nachází také pole angleRed values a torqueRed values**.

Analýza těchto dat však ukázala, že obě redukované složky obsahují výhradně nulové hodnoty – jejich přítomnost v každém záznamu nepřináší žádnou informační hodnotu, pouze zvyšuje datový objem a může komplikovat následné zpracování. Z pohledu efektivity datové infrastruktury je doporučeno, aby se **v budoucnu hodnoty angleRed a torqueRed do exportovaných JSON souborů vůbec nezapisovaly**. Tento krok by vedl k úspoře datového prostoru a k zjednodušení datové struktury nepochybně bez jakékoliv ztráty informací potřebných pro analýzu.

Výpis 6.1: Ukázka zapisování hodnot v JSON souboru (zdroj: autorka, 2026)

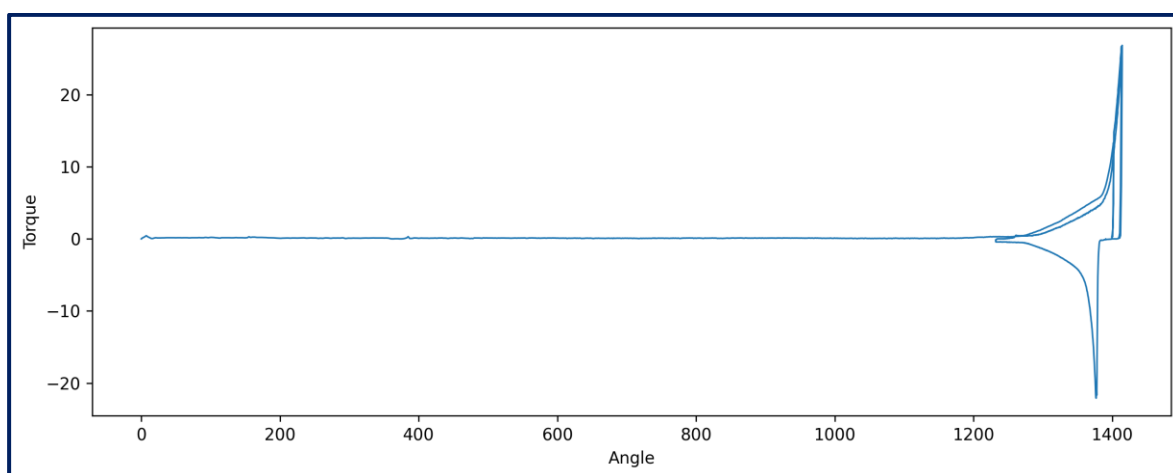
```
"graph": {
    "angle values": [150.380000, 151.170000, 153.520000,
155.680000, 158.420000, 161.170000, 164.310000, 167.640000, 169.800000, 173.520000,
177.440000, 179.600000, 183.910000, 186.070000, 190.190000, 192.340000, 196.460000,
198.820000, 202.740000, 206.660000, 210.380000, 212.540000, 216.270000, 219.990000,
224.110000, 227.830000, 231.560000, 235.290000, 237.440000, 241.360000, 245.090000,
249.010000, 251.170000, 254.890000, 258.620000, 262.740000, 266.660000, 270.380000,
274.110000, 276.460000, 280.190000, 283.910000, 288.030000, 291.750000, 295.680000,
299.400000, 303.520000, 307.240000, 311.170000, 315.090000, 319.010000, 322.730000,
326.460000, 328.810000, 332.540000, 336.460000, 340.190000, 342.340000, 346.260000,
349.990000],
    "torque values": [0.130000, 0.120000, 0.210000,
0.260000, 0.260000, 0.190000, 0.190000, 0.180000, 0.210000, 0.170000, 0.180000, 0.150000,
0.140000, 0.140000, 0.160000, 0.160000, 0.140000, 0.130000, 0.110000, 0.070000, 0.060000,
0.120000, 0.130000, 0.130000, 0.130000, 0.130000, 0.130000, 0.130000, 0.130000, 0.150000,
0.130000, 0.110000, 0.130000, 0.130000, 0.130000, 0.130000, 0.130000, 0.130000, 0.110000,
0.130000, 0.120000, 0.120000, 0.090000, 0.090000, 0.110000, 0.090000, 0.110000, 0.110000,
0.130000, 0.130000, 0.100000, 0.100000, 0.100000, 0.120000, 0.110000, 0.090000, 0.110000,
0.100000, 0.100000, 0.120000],
    "gradient values": [0.004000, 0.002000, 0.013000,
0.030000, 0.014000, -0.007000, -0.010000, -0.002000, 0.004000, 0, -0.002000, -0.005000, -
0.001000, -0.002000, 0, 0.002000, 0, -0.002000, -0.004000, -0.006000, -0.006000, 0, 0.010000,
0.007000, 0, 0, 0, 0.001000, 0.002000, 0, -0.003000, -0.001000, 0.001000, 0.001000, 0,
0.001000, 0.001000, -0.001000, -0.001000, -0.001000, 0, -0.002000, -0.002000, 0.001000, 0, 0,
0.002000, 0.004000, 0.003000, -0.002000, -0.003000, 0, 0.001000, 0, -0.001000, 0, -0.001000,
0, 0.001000],
```

```

        "torqueRed values":      [0, 0, 0, 0, 0, 0, 0, 0, 0, 0,
0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0,
0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0],
        "angleRed values": [0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0,
0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0,
0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0],
        "time values":      [0.256800, 0.258000, 0.261600,
0.264000, 0.266400, 0.268800, 0.271200, 0.273600, 0.274800, 0.277200, 0.279600, 0.280800,
0.283200, 0.284400, 0.286800, 0.288000, 0.290400, 0.291600, 0.294000, 0.296400, 0.298800,
0.300000, 0.302400, 0.304800, 0.307200, 0.309600, 0.312000, 0.314400, 0.315600, 0.318000,
0.320400, 0.322800, 0.324000, 0.326400, 0.328800, 0.331200, 0.333600, 0.336000, 0.338400,
0.339600, 0.342000, 0.344400, 0.346800, 0.349200, 0.351600, 0.354000, 0.356400, 0.358800,
0.361200, 0.363600, 0.366000, 0.368400, 0.370800, 0.372000, 0.374400, 0.376800, 0.379200,
0.380400, 0.382800, 0.385200]
    }

```

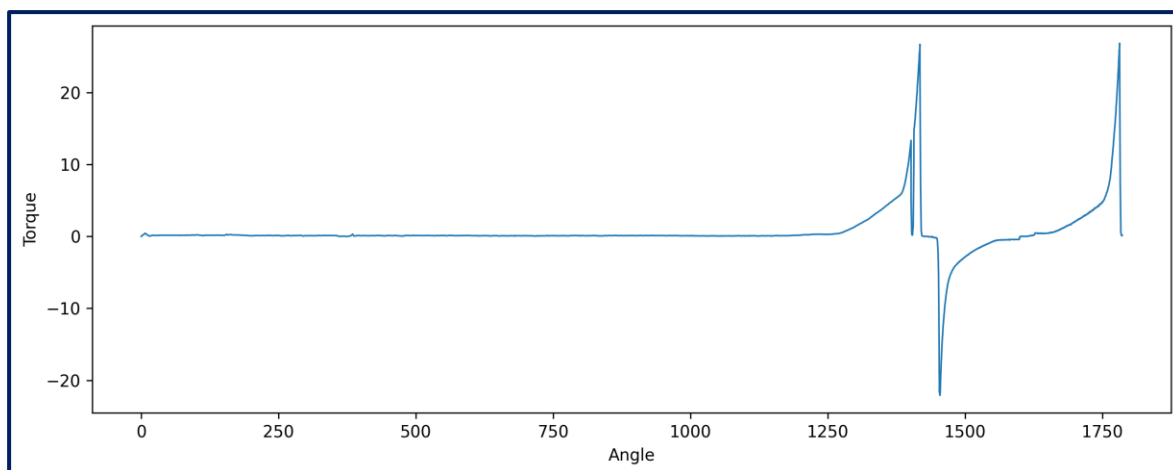
Ukázka výše představuje část původního datového záznamu ve formátu JSON, který **obsahuje průběh měření fyzikálních veličin zaznamenaných během jedné výrobní operace**. Datová struktura uložená pod klíčem graph zahrnuje několik paralelních polí, z nichž každé odpovídá jedné sledované proměnné procesu.



Obrázek 6.1: Graf standardní podoby křivky moment-úhel (zdroj: autorka, 2026)

Pole „angle values“ obsahuje úhlové hodnoty vyjádřené ve stupních, které popisují průběh pohybu nástroje v čase. S nimi koreluje pole „torque values“, jež zachycuje krouticí moment (v jednotkách N·m) a slouží jako hlavní veličina pro posouzení kvality operace. Kombinací těchto dvou veličin vzniká křivka moment–úhel, která představuje základní zdroj informací pro následnou analýzu průběhu procesu a detekci odchylek (viz **Obrázek 6.1**). Pole „gradient values“ obsahuje derivaci momentu vůči úhlu a poskytuje informace o dynamice změn zatížení. Pomocná pole „torqueRed values“ a „angleRed values“ jsou určena pro uchování zredukovaných hodnot po provedení komprese nebo zjednodušení trajektorie – v surovém záznamu jsou zatím inicializována nulami. „Time values“ pak představuje časové značky jednotlivých měření a umožňuje transformaci průběhu do časové domény, a to může být užitečné při porovnávání nebo synchronizaci více signálů.

Pro potřeby aplikace jednotlivých analytických metod byla původní data dále zpracována **s cílem vytvořit jednotný a čitelný vstupní formát pro experimentální analýzy.**



Obrázek 6.2: Graf převezené podoby křivky moment-úhel (zdroj: autorka, 2026)

Nejprve byly z každého souboru JSON **extrahovány pouze dvě klíčové veličiny – angle a torque**. Právě tyto byly totiž na základě konzultace s odborníkem z praxe vyhodnoceny jako nejvýznamnější pro posouzení průběhu dotažení a charakteristiku fyzikálního procesu. Dále byl **přidán nový sloupec angle 2, který představuje matematicky transformovanou podobu úhlové hodnoty** podle funkce navržené externím expertem. Po této transformaci byly průběhy měření po grafické stránce čitelnější a měly jednotný směr růstu (viz **Obrázek 6.2**). Pro následnou analýzu významnosti jednotlivých úseků křivky **byla přidána ke každému řádku ještě informace, o jaký segment křivky se jedná**. I přestože parametr angle ani torque nejsou časovou informací, při zachování jejich pořadí při zpracování je oba i tak lze interpretovat jako nositele informace o čase, neboť jednotlivé hodnoty odpovídají posloupnosti měření v průběhu procesu a nepřímě tak zachycují jeho časový vývoj.

Tabulka 6.1: Ukázka upraveného souboru (zdroj: autorka, 2026)

angle	torque	angle 2	segment
1148,99	26,08	1154,09	6A_Pretightening
1150,17	26,38	1155,27	6A_Pretightening
1150,17	26,59	1155,27	7A_Loosening

Výsledkem tohoto postupu je, že každý experimentální soubor ve formátu CSV obsahuje čtyři sloupce – angle, torque, angle 2 a segment, jak je možné vidět na zobrazeném příkladu tabulkové struktury (viz Tabulka 6.1). První tři sloupce představují číselné veličiny, které popisují průběh měření v oblasti momentu a úhlu, čtvrtý sloupec segment pak slouží k identifikaci konkrétní části výrobního procesu, ze které daný průběh pochází.

6.2 Komprese procesních křivek

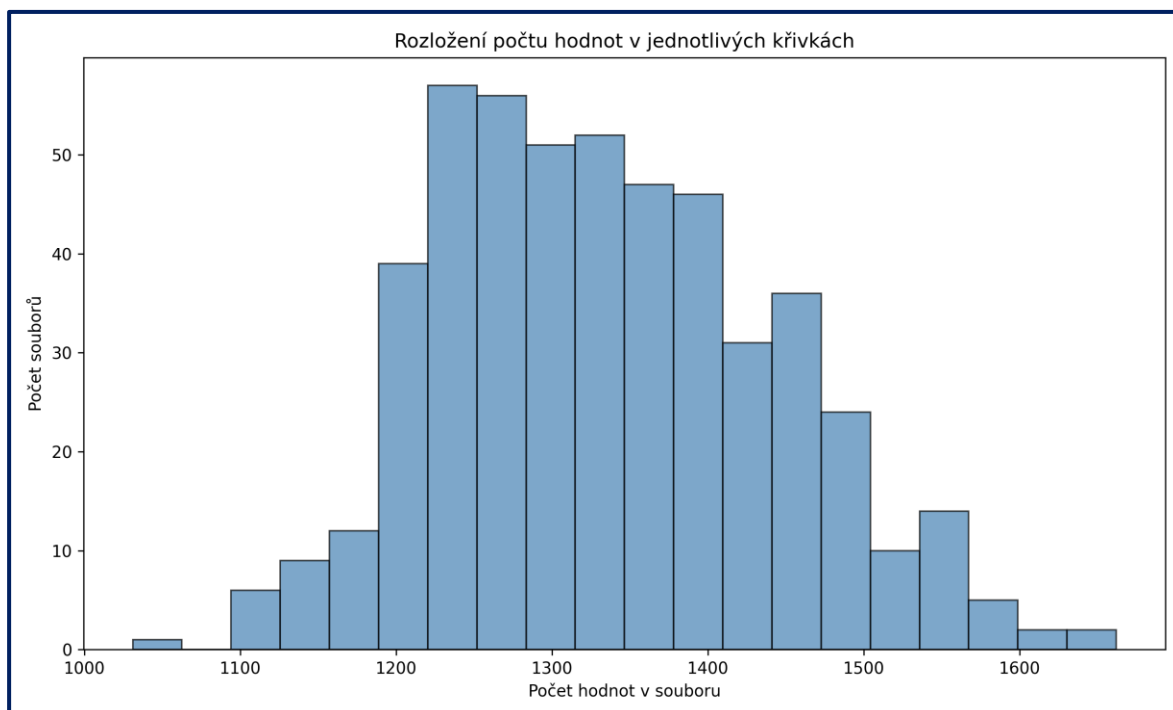
Na základě rešerše metod redukce počtu bodů na křivce byl jako kompresní nástroj **zvolen algoritmus Douglas–Peucker**, který se ukázal jako nejvhodnější pro daný kontext, neboť **umožňuje přímo redukovat datovou velikost signálu bez změny jeho reprezentace**, zachovává tvarovou informaci a současně poskytuje transparentní a snadno interpretovatelný parametr komprese v podobě prstorového prahu ϵ . V této části je popsáno, jakým způsobem byl algoritmus aplikován na dataset oříznutých výrobních křivek a jaké výsledky tato aplikace přinesla.

Aby bylo možné analyzovat vztah mezi mírou komprese a kvalitou aproximace, byla komprese provedena pro osm hodnot parametru ϵ :

- 0,05,
- 0,15,
- 0,2,

- 0,25,
- 0,3,
- 0,35,
- 0,4.

Volba konkrétních hodnot parametru ε byla provedena tak, **aby bylo možné systematicky analyzovat vliv tolerance aproximace na výslednou míru komprese a kvalitu rekonstrukce signálu.** Testované hodnoty byly zvoleny v intervalu 0,05–0,4 s krokem 0,05, což pokrylo celý relevantní rozsah parametru a zároveň poskytlo dostatečné rozlišení pro pozorování trendu změn komprese a rekonstrukční chyby. Dolní hranice intervalu byla zvolena tak, aby ještě docházelo k určité redukci počtu bodů, zatímco horní hranice představuje hodnotu, při níž by další zvyšování tolerance již vedlo k výraznějšímu zjednodušení křivky a potenciální ztrátě tvarových charakteristik signálu – interval by tedy měl poskytovat možnost analyzovat kompromis mezi mírou redukce dat a zachováním informačního obsahu procesní křivky.



Obrázek 6.3: Histogram počtu hodnot v jednotlivých souborech s křivkami (zdroj: autorka, 2026)

Na grafu rozložení počtu hodnot v jednotlivých křivkách lze **pozorovat počty hodnot na křivkách** pro lepší představu o kolik bodů opravdu lze zápis snížit.

Tabulka 6.2: Výsledky komprese s různým nastavením parametrů (zdroj: autorka, 2026)

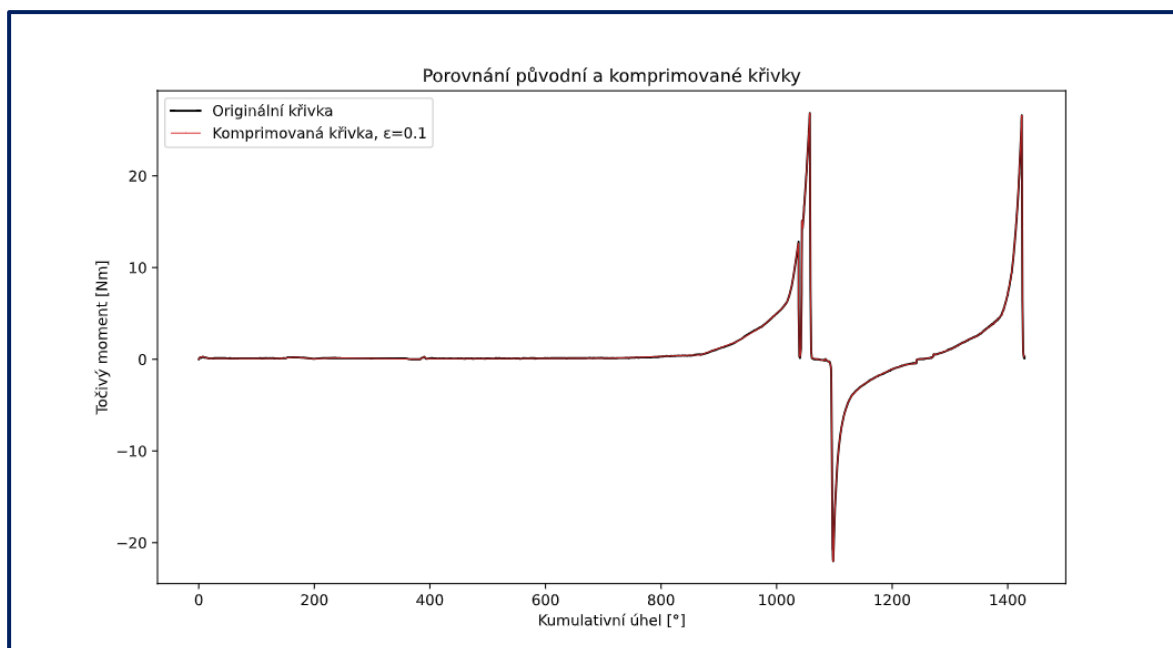
Parametr (ϵ)	Průměr výsledného počtu hodnot	Průměrná komprese (%)	Průměrná MSE
0,05	228	83%	0,0734
0,1	125	91%	0,0734
0,15	89	93%	0,0706
0,2	71	94%	0,0651
0,25	58	95%	0,0604
0,3	50	96%	0,0647
0,35	44	97%	0,0709
0,4	40	97%	0,08

Tabulka 6.2 shrnuje **vliv hodnoty parametru ϵ na úroveň komprese a přesnost rekonstrukce signálu vyjádřenou metrikou MSE (Mean Squared Error)**. Z výsledků je zřejmé, že se zvyšující hodnotou parametru ϵ se výrazně snižuje počet uchovaných bodů, což pochopitelně vede k postupnému nárůstu míry komprese. Průměrný počet bodů klesá z 228 při hodnotě $\epsilon = 0,05$ až na 40 bodů při $\epsilon = 0,4$, což odpovídá kompresi o přibližně 97 %.

Zároveň lze pozorovat, že **změny rekonstrukční chyby jsou v celém sledovaném rozsahu relativně malé** – mírný pokles hodnoty MSE v počáteční části intervalu lze interpretovat jako důsledek částečného odstranění drobných fluktuací signálu. Při vyšších hodnotách parametru ϵ však již dochází k postupnému zjednodušování průběhu křivky a rekonstrukční chyba začíná opět narůstat. Výsledky tak potvrzují existenci kompromisu mezi mírou komprese a věrností aproximace původního signálu – pro další experimenty tak byla zvolena hodnota parametru $\epsilon = 0,25$, při níž je dosaženo komprese okolo 95 % při současně nejmenší průměrné hodnotě rekonstrukční chyby v rámci testovaného intervalu.

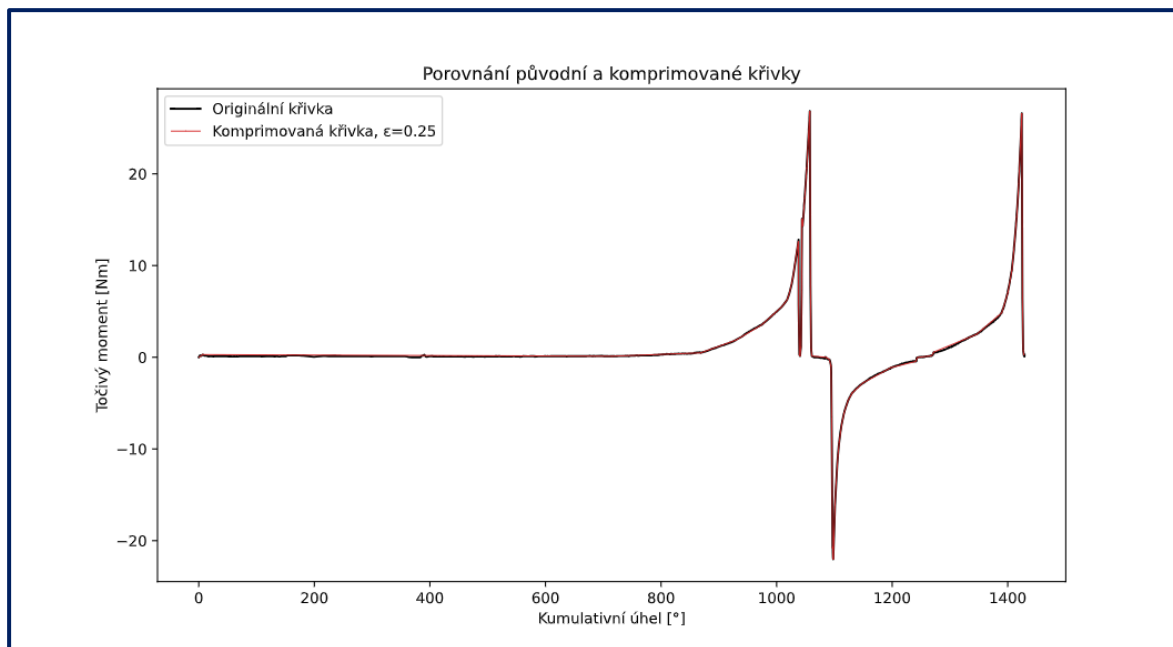
Rozhodující **kritérium pro optimální volbu $\epsilon = 0,25$** není samotná minimální MSE, ale **kombinace tří faktorů**: nejmenší MSE v testovaném intervalu, potvrzení expertního posouzení dvěma procesními specialisty a dosažení komprese 95 %, která představuje prakticky smysluplnou hranici pro archivaci. Volba byla dále ověřena vizuální inspekcí na reálných defektních křivkách.

Pro ilustraci **vlivu parametru ϵ na tvar komprimované křivky jsou na obrázcích nejprve uvedeny příklady aproximace procesní křivky při nejmenší, zvolené optimální a nejvyšší úrovni tolerance**.



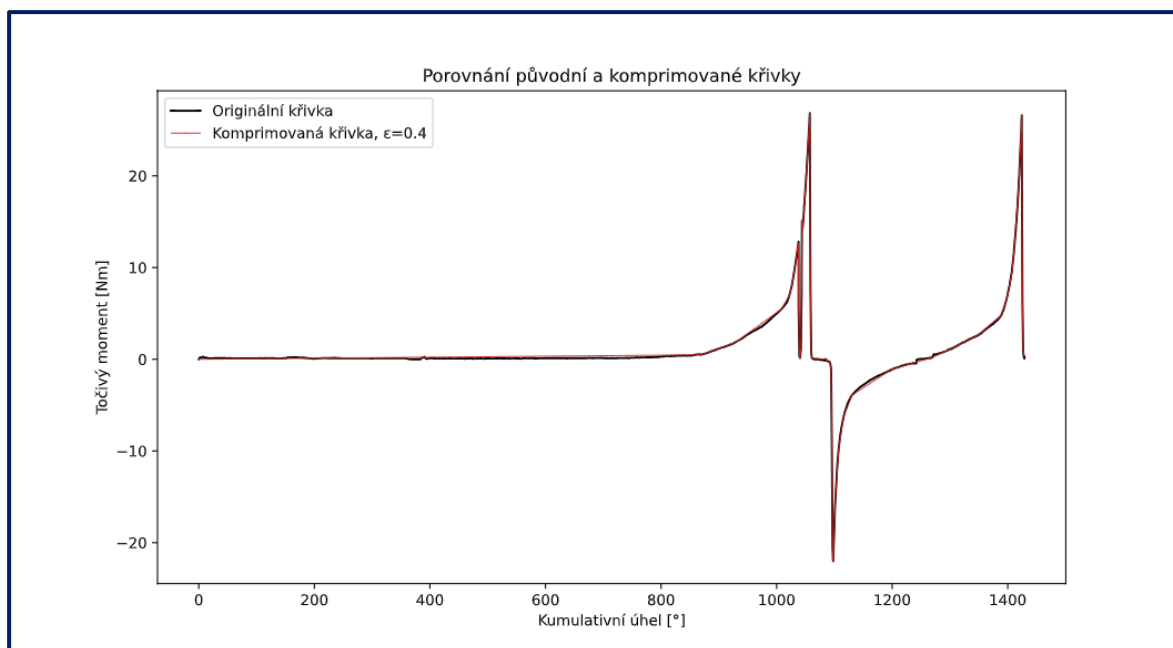
Obrázek 6.4: Porovnání původní a komprimované procesní křivky při parametru $\epsilon = 0,05$ (zdroj: autorka, 2026)

Obrázek s grafem zobrazuje **porovnání původní křivky s její komprimovanou reprezentací** při nejnižším nastavení $\epsilon = 0,05$. Na obrázku lze vidět, že při této hodnotě dochází pouze k mírné redukci počtu bodů a komprimovaná křivka téměř přesně kopíruje průběh původního signálu – zachovány zůstávají všechny významné tvarové charakteristiky a vizuální rozdíly mezi oběma průběhy jsou minimální: výsledná aproximace zůstává velmi blízká původním datům.

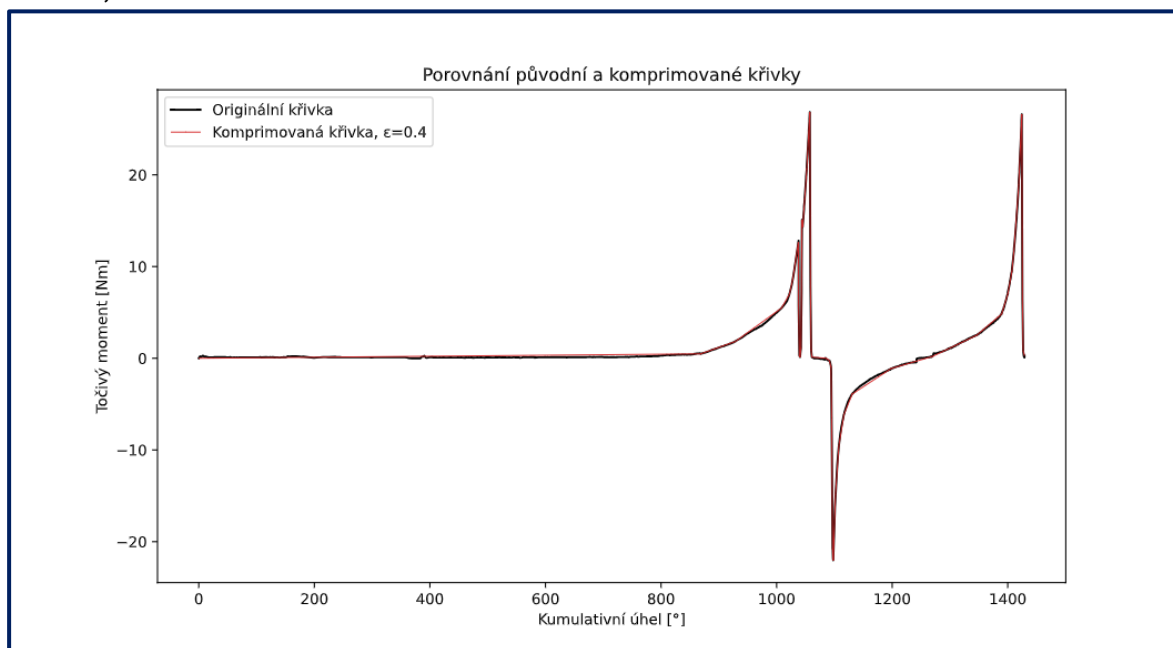


Obrázek 6.5: Porovnání původní a komprimované procesní křivky při parametru $\epsilon = 0,25$ (zdroj: autorka, 2026)

Na obrázku (viz **Obrázek 6.5**) je zachyceno **porovnání původní a komprimované křivky při optimální hodnotě parametru $\epsilon = 0,25$** . V tomto případě již dochází k výraznější redukci počtu bodů a k částečnému zjednodušení průběhu signálu. Přesto komprimovaná křivka nadále zachovává hlavní trend průběhu i charakteristické body procesu.



Obrázek 6.6: Porovnání původní a komprimované procesní křivky při parametru $\varepsilon = 0,4$ (zdroj: autorka, 2026)



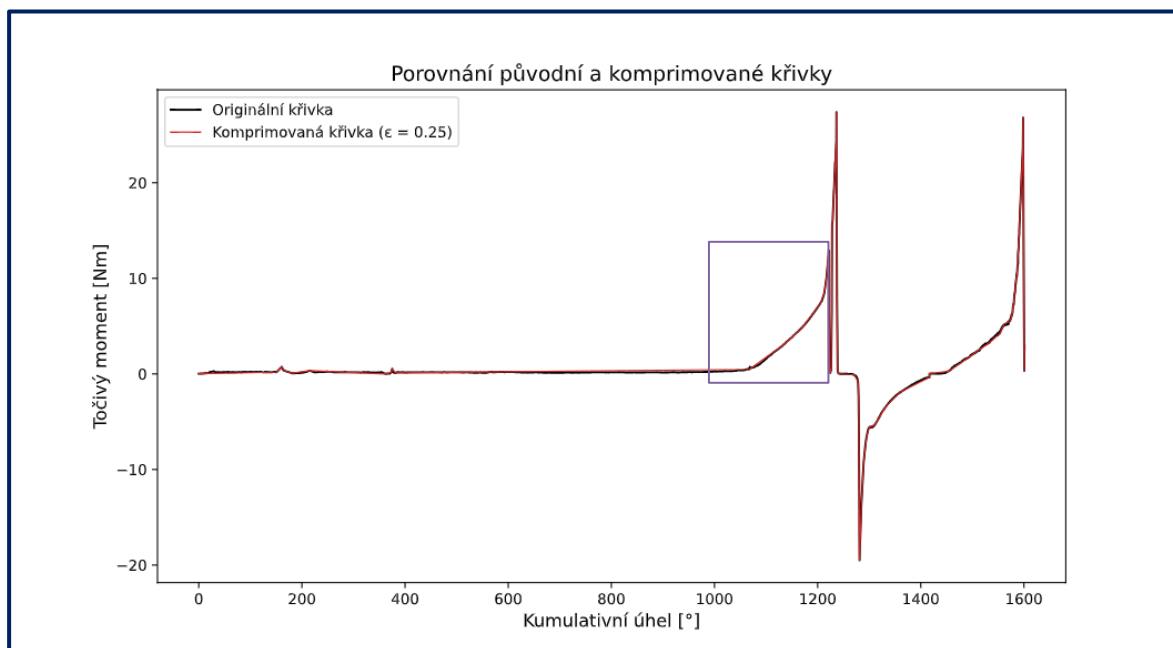
Obrázek 6.6 zobrazuje **porovnání původní procesní křivky s její komprimovanou reprezentací vytvořenou algoritmem Douglas–Peucker při nastavení parametru $\varepsilon = 0,4$** . Nejvyšší testovaná hodnota, a tedy i nejvyšší možná hodnota komprese stále zachovává poměrně detailní tvar průběhu procesu a hlavní charakteristické body signálu, tedy výrazné vrcholy a prudké změny momentu. V porovnání s nižšími hodnotami parametru ε každopádně dochází ke zjednodušení průběhu a k eliminaci drobných lokálních fluktuací, nicméně výsledný průběh nadále věrně reprezentuje hlavní trend křivky.

Aby bylo možné posoudit, zda je zvolená metoda komprese použitelná v reálném výrobním prostředí, nestačí hodnotit pouze obecné metriky – z pohledu kvality výroby je nejdůležitější, zda komprimovaná křivka stále obsahuje dostatek informace pro spolehlivé odhalení nežádoucího chování procesu, tedy zejména anomálních kusů, které by u zákazníka vedly k reklamaci.

Pro experimentální ověření byly proto využity záznamy křivek z reálného provozu, které reprezentují takové kusy, u kterých lze přepokládat jejich vadnost, ačkoliv je tradiční metody označily za OK –

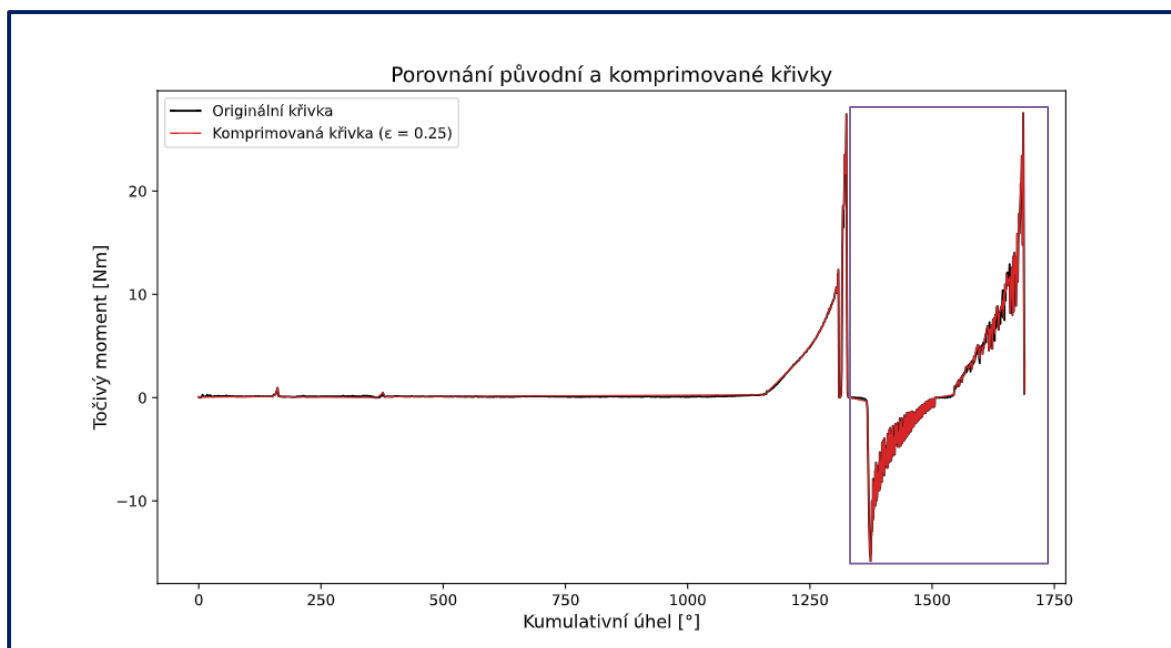
konkrétně takové, které by byly s velmi vysokou pravděpodobností reklamovány, pokud by se dostaly k zákazníkovi.

Doplňková konzultace se dvěma experty na daný výrobní proces potvrdila, že i **zvažovaná hodnota parametru $\epsilon = 0,25$ poskytuje pro účely archivace plně dostačující úroveň** zachování tvaru křivky. Z tohoto důvodu byly následné experimenty prováděny právě s touto hodnotou epsilon.



Obrázek 6.7: Testovací porovnání původní a komprimované procesní křivky č. 1 při parametru $\epsilon = 0,25$ (zdroj: autorka, 2026)

Na obrázku (viz **Obrázek 6.7**) je patrné, že i **při více než 94% kompresi křivky zůstává průběh signálu prakticky totožný s originálem při hodnotě MSE 0,059**. Komprimovaná křivka věrně zachovává jak globální tvar, tak i lokální extrémy a prudké změny, včetně oblasti mezi body 1000 a 1200 na ose x, kde by se v případě potřeby prováděla zpětná analýza chyby. Z pohledu následné diagnostiky tedy nedochází ke ztrátě informací, které jsou nepostradatelné pro identifikaci potenciálně vadných kusů.



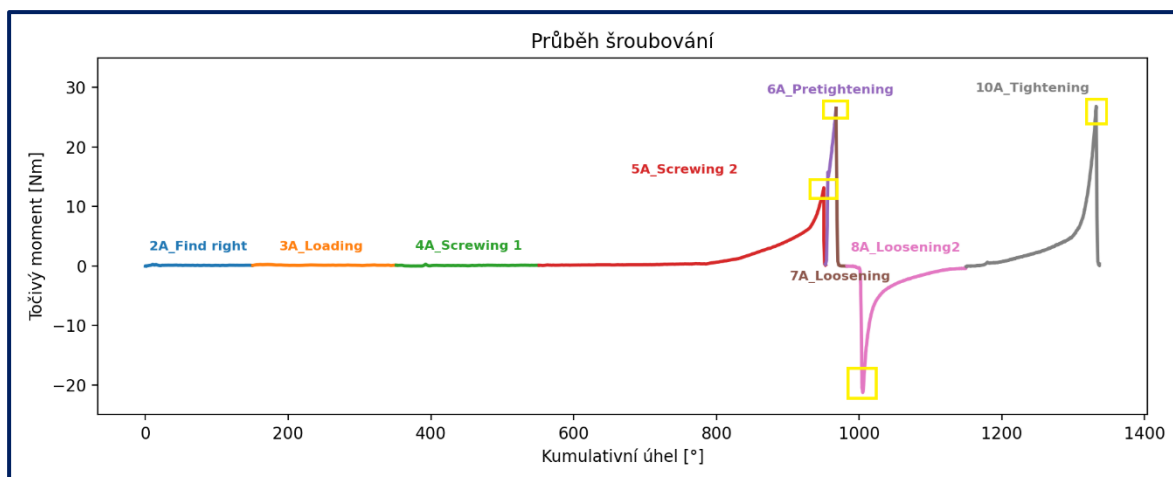
Obrázek 6.8: Testovací porovnání původní a komprimované procesní křivky č. 2 při parametru $\epsilon = 0,25$ (zdroj: autorka, 2026)

Pro lepší otestování metody byl do testovací množiny křivek zařazen i na první pohled defektní díl, který však tradičními způsoby kontroly projde jako OK – i přes velmi nestandardní průběh v žlutě vyznačené oblasti se ho podařilo s nastavením parametru $\epsilon = 0,25$ komprimovat na 69 %, při kterých hodnota MSE nabyla hodnoty 0,0824.

Díky provedenému testování na reálných výrobních datech lze jednoznačně **potvrdit, že algoritmus Douglas–Peucker poskytuje plně funkční a spolehlivý postup** pro kompresi křivkových dat určených k dlouhodobé archivaci a z praktického hlediska tak představuje vhodný nástroj pro optimalizaci datových úložišť, zejména v prostředích, kde je generováno velké množství sensorických záznamů a kde je zároveň vyžadováno zachování věrného tvaru křivky pro účely pozdějšího vyhodnocení nebo auditních procesů.

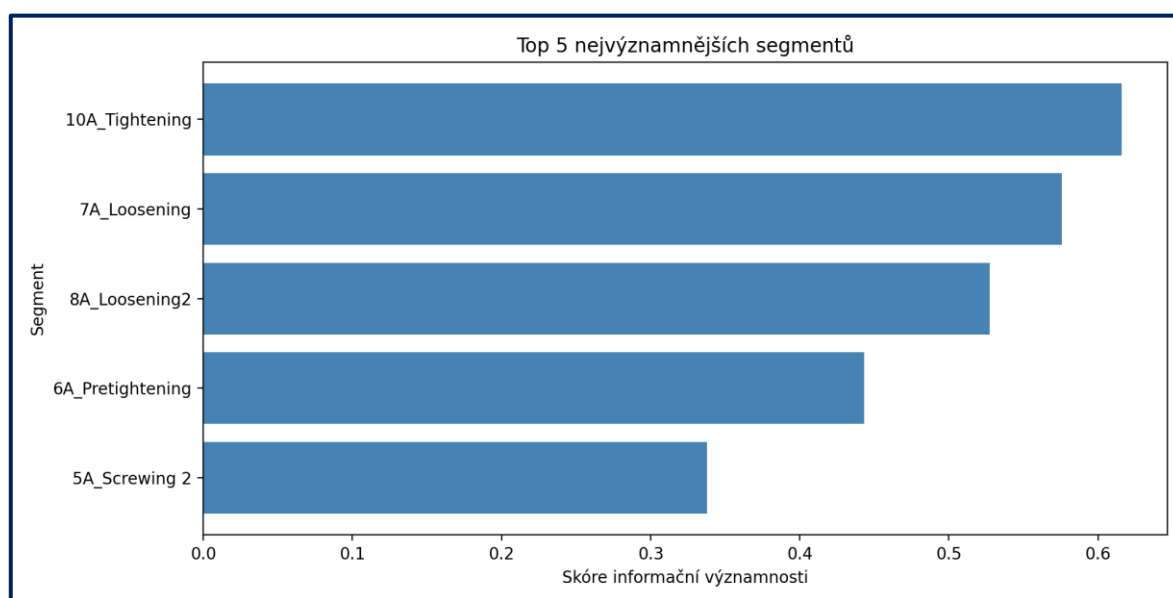
6.3 Významnost jednotlivých úseků křivky

Cílem této fáze bylo **najít ty části křivky moment–úhel, které nesou největší množství diagnostických informací** a jsou proto nejdůležitější pro sledování případných odchylek od standardního chování procesu. Zároveň lze díky této analýze potvrdit nebo vyvrátit správnost umístění kontrolních oken pro automatické vyhodnocování OK a NOK kusů (viz **Obrázek 6.9**), která jsou na obrázku vyznačena žlutým ohraničením.



Obrázek 6.9: Graf průběhu křivky moment-úhel s vyznačenými kontrolními okny a jednotlivými segmenty (zdroj: autorka, 2026)

Pro analytické účely byla implementována metoda pro vyhodnocení informační významnosti jednotlivých úseků křivky – **každý segment** (viz **Obrázek 6.9**) **byl vyhodnocen samostatně podle čtyř metrik: směrodatné odchylky momentu, rozsahu hodnot momentu, plochy pod křivkou a variability sklonu**. Pro zajištění vzájemné srovnatelnosti byly všechny metriky převedeny do jednotného měřítka v intervalu 0–1 a následně bylo vypočteno souhrnné skóre významnosti jako vážený průměr jejich hodnot. Segmenty vykazující výrazné změny momentu, větší amplitudu a rychlejší přechody dosahovaly vyšších hodnot skóre, což signalizuje jejich vyšší informační přínos. Naopak části průběhu s ustáleným charakterem byly vyhodnoceny jako méně významné.



Obrázek 6.10: Graf významnosti segmentů (zdroj: autorka, 2026)

Obrázek 6.10 zobrazuje **pět segmentů s nejvyšším dosaženým skóre informační významnosti**. Z grafu je patrné, že dominantními oblastmi z hlediska přínosu pro kontrolu kvality jsou segmenty 10A, 7A a 8A, kvůli výrazné dynamice a rychlým změnám momentu. Výsledky potvrzují, že největší množství diagnostických informací je soustředěno právě v aktivních fázích procesu, zatímco stabilní části průběhu mají na celkové vyhodnocení menší vliv; jsou v souladu s expertním hodnocením, podle něhož byla tzv. okénková analýza zaměřena na nejvyšší a nejnižší části křivky, konkrétně v segmentech 10A_Tightening, 8A_Loosening2 a mezi kroky 7A_Loosening a 6A_Pretightening. Je nutné zmínit, že segment 5A_Screwing2 byl v oblasti růstu experty označen za problémový ve smyslu naprosto nejčastějšího výskytu anomálií na křivce, a tedy také koresponduje se zjištěním významnosti jednotlivých segmentů.

6.4 Detekce anomálií na křivkách

Detekce anomálií na křivkových datech byla realizována **pomocí konvolučního autoenkodéru**, který byl navržen pro rekonstrukci běžných průběhů momentové křivky. Základní princip metody spočívá v tom, že **model je učen pouze na křivkách odpovídajících standardnímu chování procesu**, a následně je **schopen tyto průběhy rekonstruovat s nízkou chybou** (Meiners et al., 2020). Pokud je na vstup předložena křivka s nestandardním průběhem, tedy potenciálně anomální záznam, rekonstrukční chyba se zpravidla zvýší. Právě velikost této chyby byla v dalším kroku využita jako kritérium pro rozlišení mezi standardními a anomálními kusy.

Pro experiment byly použity výhradně **hodnoty torque seřazené v čase**, tedy průběhy krouticího momentu, neboť technologický význam této veličiny fakticky poskytuje hlavní zdroj informace o průběhu dotažení.

Protože délka jednotlivých průběhů se výrazně lišila (viz **Obrázek 6.3**, typicky mezi 1 000 až 1 600 body), bylo pro potřeby detekce anomálií nezbytné sjednotit datový rozsah všech měření. Původně byla použita interpolace na 1 000 bodů, která měla zajistit zachování charakteru křivky při zmenšení datového objemu; tato metoda bohužel vykazovala nedostatečnou přesnost zejména v oblastech s prudkými změnami momentu. Po konzultaci s procesními odborníky bylo rozhodnuto o **použití alternativního přístupu – ořezání průběhů odzadu na délku 1 000 bodů**. Způsob byl zvolen proto, že úvodní fáze křivky, odpovídající počátečnímu ustavení systému, se ukázala jako informačně málo významná a z hlediska analýzy procesu jen zřídka přinášela diagnosticky relevantní údaje, zároveň obsahovala poměrně různící se počty bodů vzhledem k tomu, že při tomto měření záleží na úhlu zasažení kusu do výrobní linky. Každý záznam tedy nadále vystupoval jako jednorozměrný číselný vektor o stejné délce.

Trénovací množina byla tvořena 502 soubory, které reprezentovaly pouze standardní, tedy bezvadné průběhy procesu. Pro potřeby evaluace byly zaslány dvě složky pro OK a anomální křivky, z nichž každá obsahovala 30 křivek – je nutné zmínit, že evaluační složka s anomálními křivkami obsahovala jak křivky vyhodnocené jako OK, tak ty, které tradiční metody označily za NOK.

Před vstupem do modelu byla **data škálována pomocí Min–Max normalizace**. Normalizační transformace byla odvozena z trénovacích dat a následně aplikována také na evaluační křivky. Výsledné hodnoty tak byly převedeny do intervalu od 0 do 1, což je u autoenkodérů pro detekci anomálií v časových řadách nejen doporučený postup (TensorFlow, 2024), ale současně koresponduje s využitím sigmoidální aktivační funkce na výstupu (Goodfellow et al., 2016a). Po škálování byla data doplněna o kanálovou dimenzi tak, aby odpovídala vstupnímu formátu konvoluční sítě ve tvaru ((1000, 1)), protože jednorozměrná konvoluční síť očekává vstup ve formě sekvence doplněné o informaci o počtu vstupních kanálů (Goodfellow et al., 2016b).

Samotný model byl implementován jako **sekvenční architektura typu CNN autoenkodér**, jehož cílem je rekonstruovat vstupní signál na základě naučené vnitřní reprezentace (Goodfellow et al., 2016a). Vstupní vrstva přijímala jednorozměrnou křivku délky 1 000 bodů. Enkodérová část modelu byla tvořena třemi konvolučními vrstvami s 64, 32 a 16 filtry a s velikostmi kernelu 7, 5 a 3. Využití více vrstev tvoří tzv. hluboký autoenkodér, který oproti mělkým strukturám dokáže data mnohem lépe komprimovat a potenciálně snižovat výpočetní náročnost pro reprezentaci některých funkcí (Goodfellow et al., 2016a). Navíc platí, že využití konvolučních kernelů zajišťuje řídké propojení, což ve své podstatě poskytuje ekvivalenci vůči posunutí – síť díky tomu dokáže zachytit stejné charakteristiky nezávisle na tom, v jakém čase se v křivce objeví (Goodfellow et al., 2016b). Redukce dimenze skrze pooling navíc vytváří tzv. **neúplný autoenkodér, jehož kapacita je záměrně omezena** tak, aby nedokázal signál dokonale zkopírovat, ale naopak **byl donucen zachytit ty nejvýraznější vlastnosti trénovacích dat** (Goodfellow et al., 2016a).

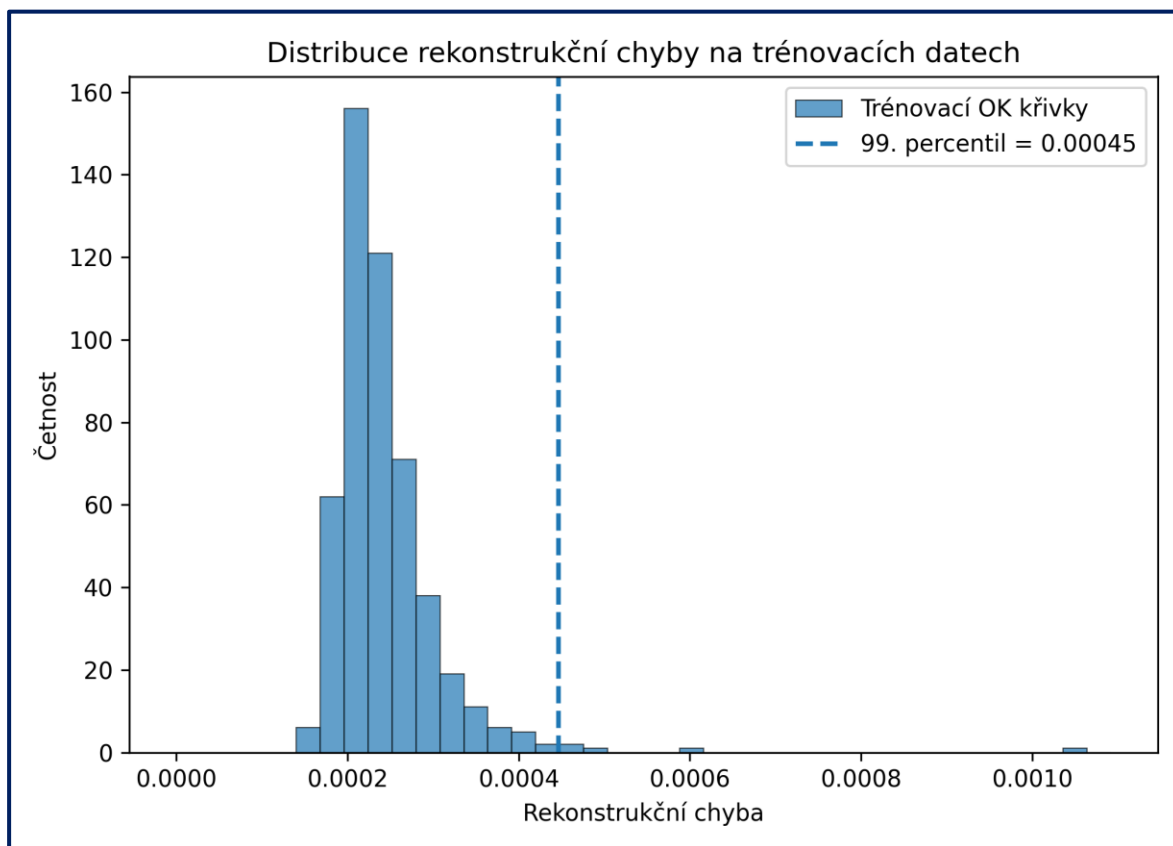
Za prvními dvěma konvolučními vrstvami enkodéru byly zařazeny vrstvy MaxPooling1D, které postupně redukovaly rozměr vnitřní reprezentace a umožnily modelu **soustředit se na hlavní tvarové charakteristiky signálu**. Operace **sdužování (tedy pooling)** nahrazuje výstupy sítě souhrnnou statistikou z blízkého okolí, čímž činí reprezentaci invariantní vůči malým lokálním posunům ve vstupu. Současně efektivně zmenšuje prostorovou dimenzi dat a snižuje tím výpočetní zátěž v následujících vrstvách (Goodfellow et al., 2016b). **Dekodérová část následně rekonstruovala původní průběh** pomocí dvou operací UpSampling1D prokládaných konvolučními vrstvami, přičemž výstupní vrstva s jedním filtrem a sigmoidální aktivační funkcí vracela rekonstruovanou křivku ve stejném formátu jako vstup.

Model byl optimalizován **pomocí algoritmu Adam a jako ztrátová funkce byla použita MSE** – to teoreticky odpovídá standardnímu trénování autoenkodérů, které se běžně optimalizují pomocí variant dávkového gradientního sestupu s využitím zpětného šíření chyby, přičemž cílem ztrátové funkce je penalizovat model za jakoukoliv odlišnost rekonstrukce od původního vstupu (Goodfellow et al., 2016a). Trénování probíhalo maximálně po dobu 300 epoch s velikostí dávky 32 a s validačním podílem 10 % z trénovacích dat. Pro omezení přeučení byla současně využita metoda early stopping, která ukončila učení ve chvíli, kdy se validační ztráta po dobu dvaceti po sobě jdoucích epoch dále nezlepšovala – v takovém případě by se obnovily váhy z epochy s nejnižší hodnotou validační ztráty.

Po dokončení trénování byla **pro každou evaluační křivku vypočtena její rekonstrukce**, což se v podstatě dá označit za samotné jádro detekce anomálií pomocí autoenkodérů – vychází totiž z hypotézy, že pokud je model trénován primárně na normálních datech, abnormální (anomální) průběhy nedokáže přesně zkopírovat a jejich rekonstrukční chyba tak bude prokazatelně vyšší (TensorFlow, 2024). Následně byla **stanovena rekonstrukční chyba**, která nebyla počítána jako prostá globální MSE přes celý průběh, ale jako vážená rekonstrukční chyba. Ačkoliv se v základních úlohách často využívá průměrná globální chyba porovnávána s pevným prahem, správný přístup a volba strategie pro určení anomálií vždy závisí na specifikách konkrétní datové sady (TensorFlow, 2024). Z tohoto důvodu byla navržena metodika zohledňující tvarové vlastnosti signálu. Každá křivka byla rozdělena na dvě části: první úsek zahrnoval prvních 250 bodů a druhý úsek zbývajících 750 bodů. Pro obě části byla samostatně spočtena střední kvadratická chyba rekonstrukce a celková hodnota byla určena jako vážený součet

$$MSE_j = w_1 \cdot MSE_{j,1} + w_2 \cdot MSE_{j,2}, \quad (6.8)$$

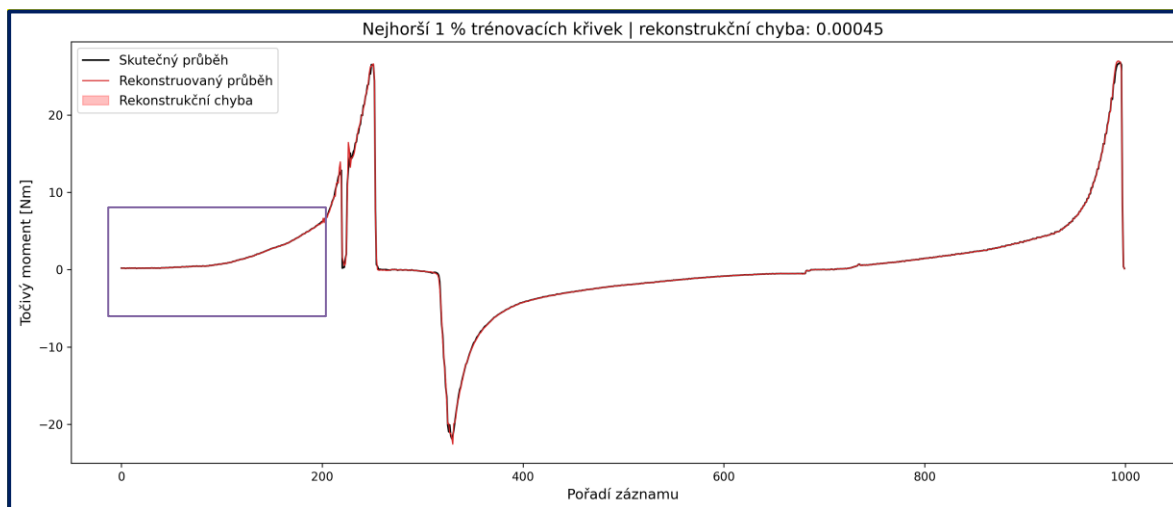
kde $MSE_{j,1}$ představuje chybu rekonstrukce první části j -té křivky, $MSE_{j,2}$ chybu rekonstrukce druhé části a váhy byly nastaveny na $w_1 = 0,7$ a $w_2 = 0,3$. Vyšší váha byla přiřazena první části průběhu, protože se při expertním ověřování ukázala jako výrazně významnější pro rozlišení běžného a nestandardního chování.



Obrázek 6.11: Distribuce rekonstrukční chyby na trénovacích datech (zdroj: autorka, 2026)

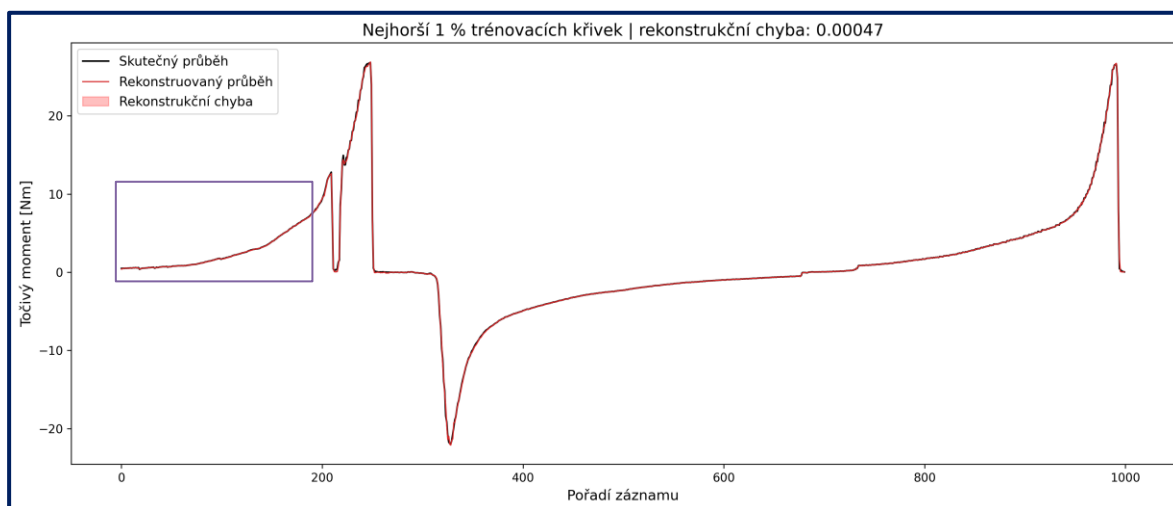
Pro převod na binární klasifikaci OK/NOK byl následně **stanoven rozhodovací práh na základě rozdělení rekonstrukční chyby na trénovacích datech**, konkrétně jako 99. percentil rekonstrukční

chyby trénovacích OK křivek. Křivky s rekonstrukční chybou vyšší než tento práh byly klasifikovány jako anomální.



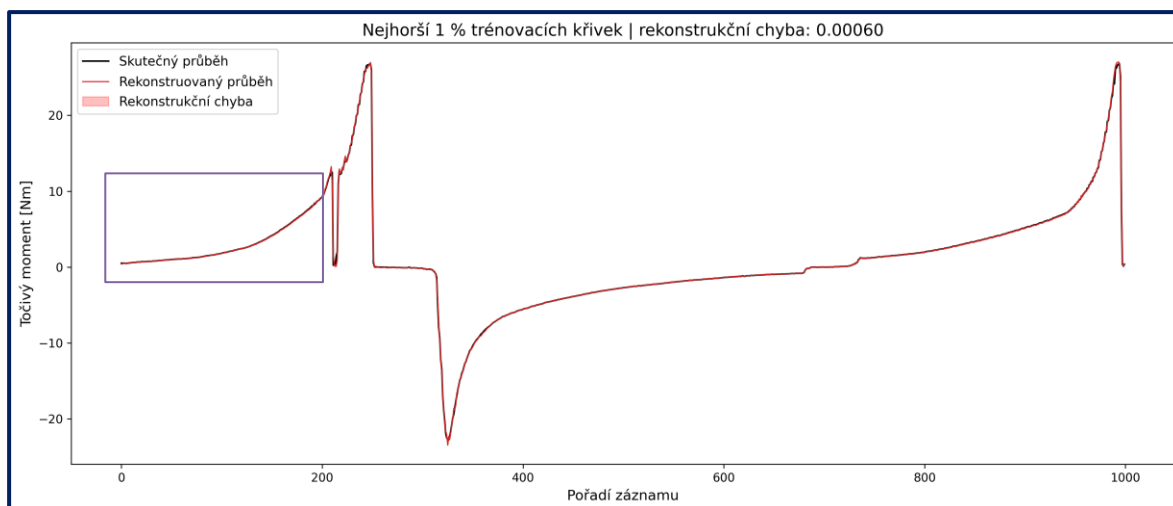
Obrázek 6.12: První křivka přesahující threshold (zdroj: autorka, 2026)

Na zobrazené křivce lze pozorovat **rozporuplný nárůst mezi body 0–200**. Popisovaná oblast je na této křivce, a i na dalších zobrazovaných zvýrazněna žlutým ohraničením.



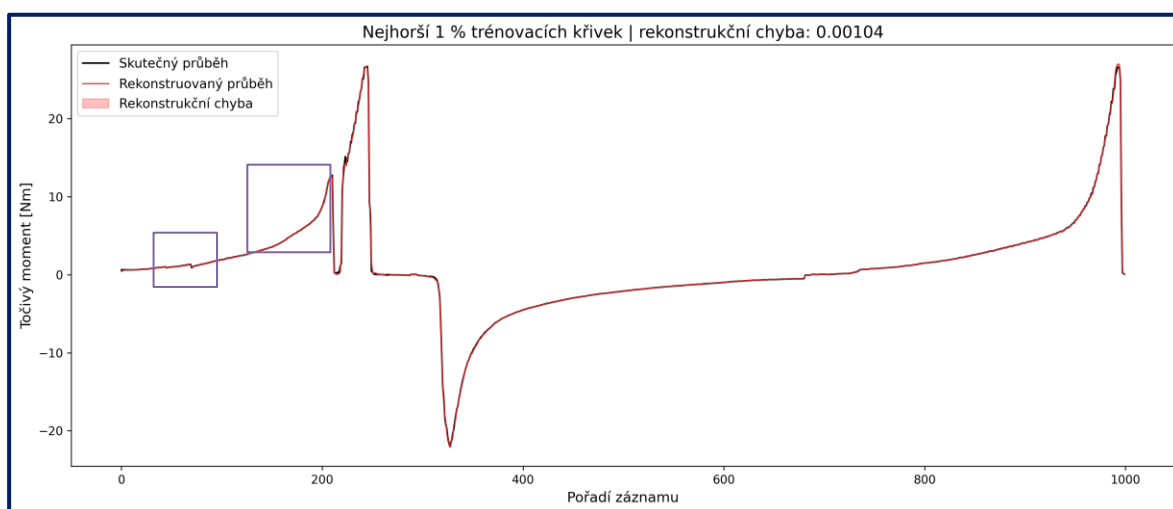
Obrázek 6.13: Druhá křivka přesahující threshold (zdroj: autorka, 2026)

Druhá **rekonstruovaná křivka přesahující 99. percentil** opět ukazuje rozporuplný nárůst na svém počátku.



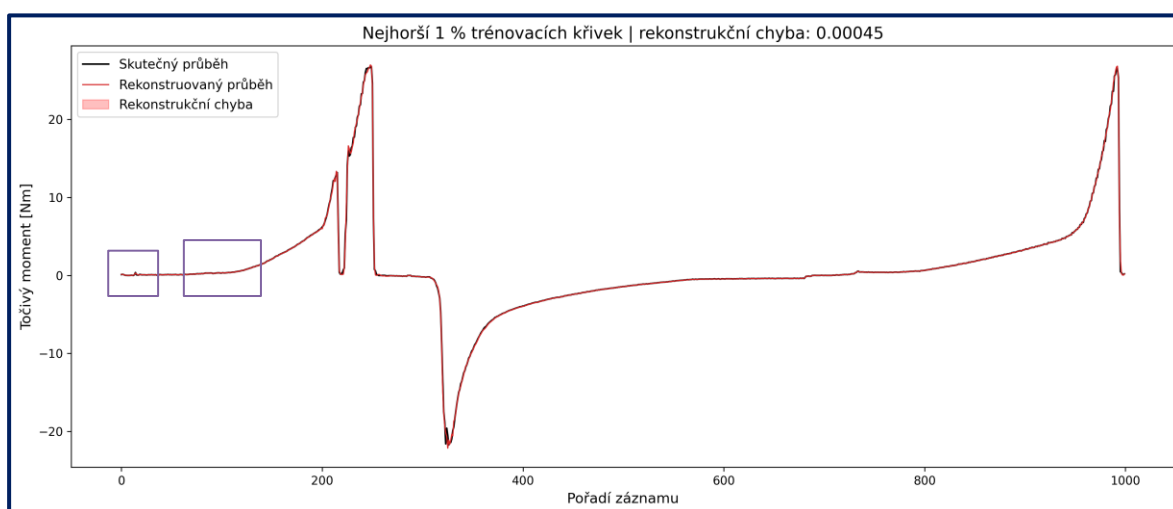
Obrázek 6.14: Třetí křivka přesahující threshold (zdroj: autorka, 2026)

Na třetí křivce je stejný problém jako na první a druhé.



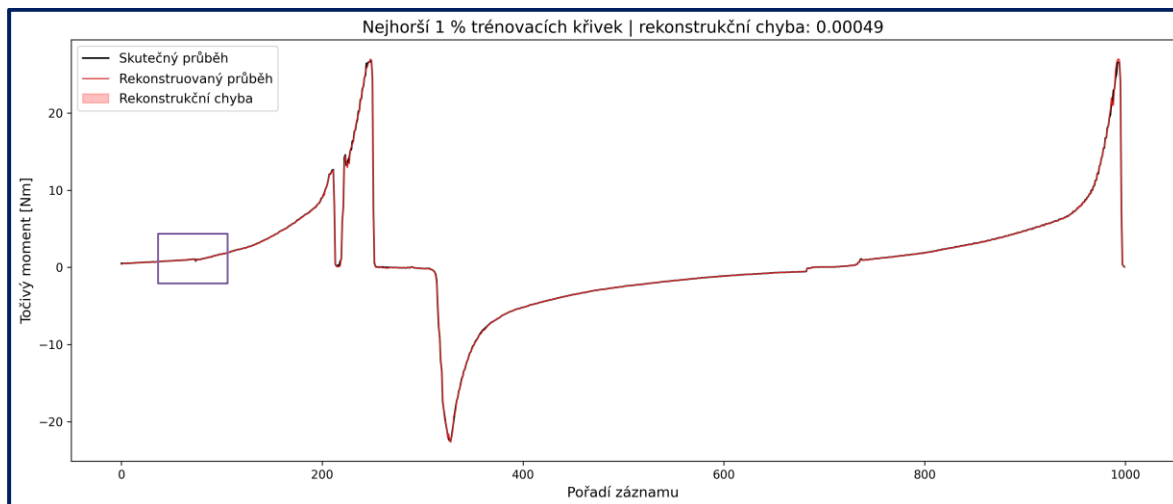
Obrázek 6.15: Čtvrtá křivka přesahující threshold (zdroj: autorka, 2026)

Na čtvrté křivce je mezi 0–100 bodem patrný zub a nestandardní nárůst okolo 200. bodu.



Obrázek 6.16: Pátá křivka přesahující threshold (zdroj: autorka, 2026)

Pátá křivka má v počátku malý zub a okolo 100. bodu viditelnou malou vlnku.



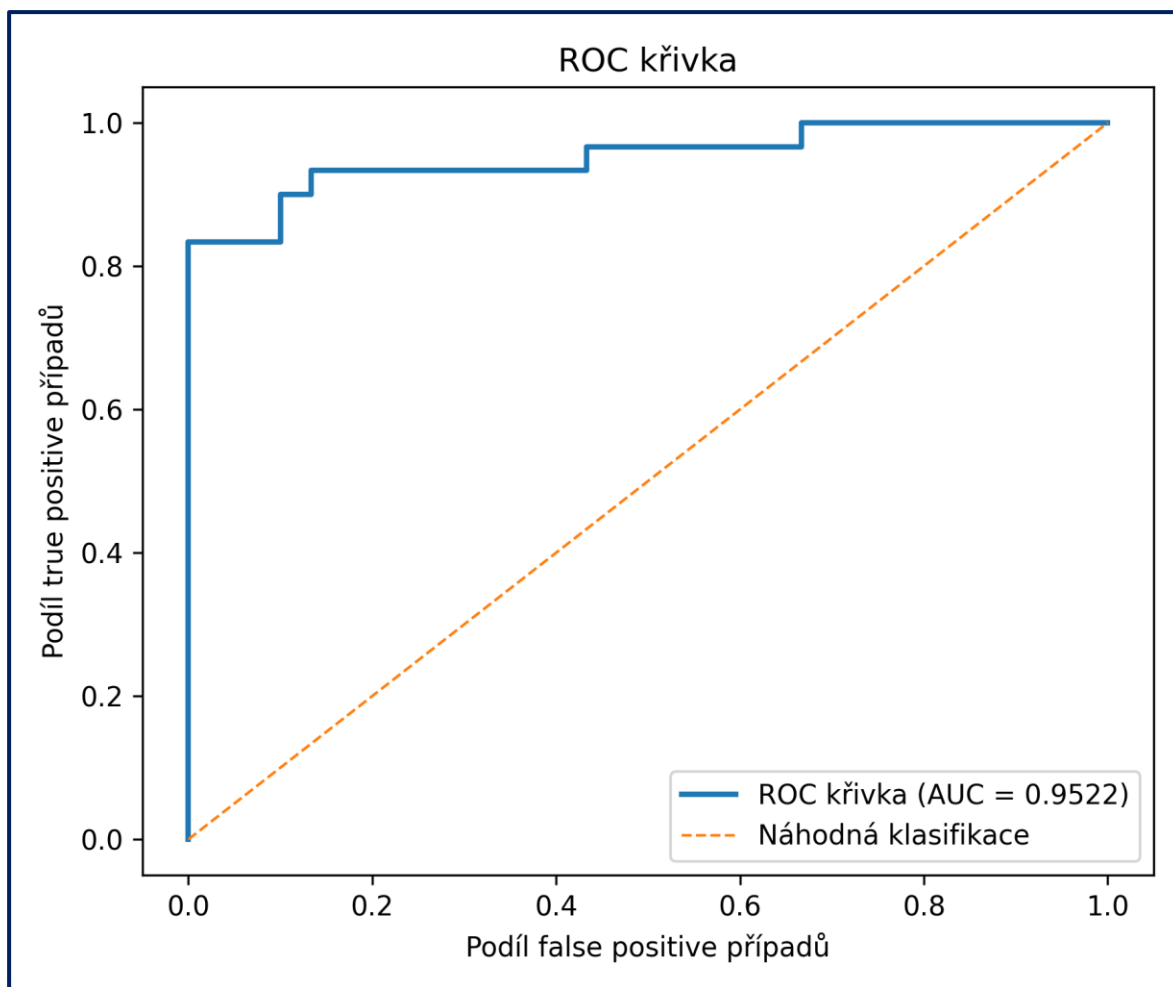
Obrázek 6.17: Šestá křivka přesahující threshold (zdroj: autorka, 2026)

Na šesté křivce je, stejně jako na čtvrté, viditelný o něco menší zub.

Z uvedených ukázek je vidět, že **křivky překračující stanovený práh mají v průběhu určité lokální odchylky, nejčastěji na začátku měření**. Typicky jde o různé náhlé nárůsty, „zuby“ nebo drobné vlnění, které se modelu nedaří dobře zrekonstruovat, a proto u nich vychází vyšší rekonstrukční chyba; zároveň je ale potřeba říci, že ne ve všech případech jde o jednoznačně vadné kusy – spíše se jedná o hraniční průběhy, které se už částečně odlišují od běžného chování.

Kvalita klasifikace byla hodnocena pomocí standardních metrik, konkrétně accuracy, precision, recall a F1-score, a zároveň byla sestavena také konfuzní matice s vidinou podrobného posouzení zastoupení správně a chybně klasifikovaných případů.

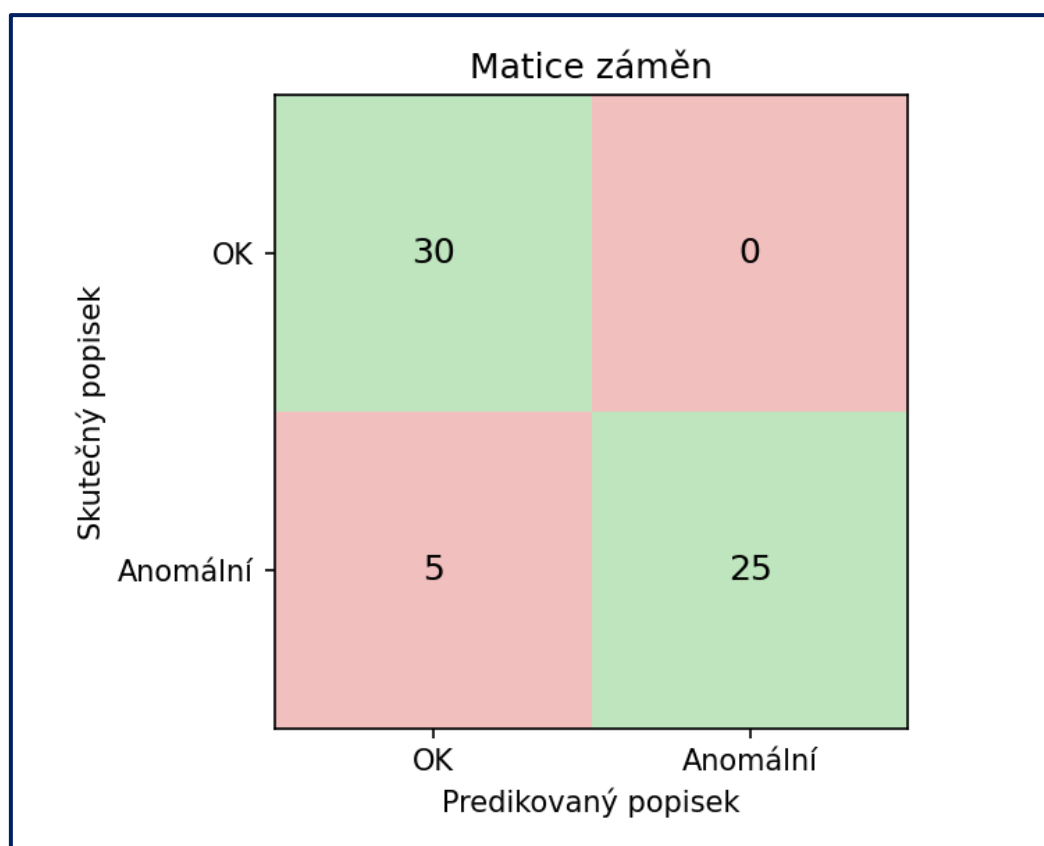
Pro lepší interpretaci výsledků byly všechny evaluační křivky po klasifikaci také graficky exportovány – každý výstupní graf obsahoval původní torque křivku, její rekonstrukci vytvořenou autoenkodérem a červeně zvýrazněnou oblast mezi oběma průběhy, která vizuálně reprezentovala rekonstrukční chybu. Grafy byly ukládány do samostatných podsložek podle výsledku klasifikace, tedy zvlášť pro správně a chybně vyhodnocené případy (jak true positive a negative, tak false positive a negative). Z teoretického hlediska tyto vizuální odchylky **přímo ukazují místa, kde vstupní signál opouští prostor normálních dat**. Autoenkodér si totiž zachovává schopnost reprezentovat pouze ty variace, které jsou nezbytné pro úspěšnou rekonstrukci normálních trénovacích vzorků, a cokoliv mimo tuto oblast nedokáže přesně zkopírovat (Goodfellow et al., 2016a).



Obrázek 6.18: ROC křivka modelu pro detekci anomálií na momentových křivkách (zdroj: autorka, 2026)

Prvním krokem hodnocení výkonnosti navrženého modelu byla **analýza ROC křivky, která zachycuje vztah mezi podílem správně detekovaných anomálií a podílem falešně pozitivních případů**, tedy normálních křivek označených za anomálie při různých hodnotách rozhodovacího prahu.

Z grafu (viz **Obrázek 6.18**) je patrné, že křivka se nachází výrazně nad diagonálou reprezentující náhodnou klasifikaci, a to indikuje dobrou schopnost modelu rozlišovat mezi standardními a anomálními průběhy momentové křivky. **Celková diskriminační schopnost modelu byla kvantifikována pomocí plochy pod ROC křivkou**, která v tomto experimentu dosáhla hodnoty 0,9522, tedy výrazně vyšší, než dosahuje náhodná klasifikace s hodnotou 0,5.



Obrázek 6.19: Matice záměn (zdroj: autorka, 2026)

Matice záměn (viz **Obrázek 6.19**) přehledně zachycuje **vztah mezi skutečným a predikovaným označením jednotlivých evaluačních křivek a umožňuje rozlišit čtyři základní typy výsledků klasifikace**: správně identifikované normální průběhy (true negative), správně detekované anomálie (true positive), falešně pozitivní případy (false positive) a falešně negativní případy (false negative).

Z matice záměn vyplývá, že **model správně klasifikoval všech 30 normálních křivek**, které byly označeny jako OK, v případě anomálních křivek bylo správně detekováno 25 průběhů, zatímco 5 anomálních křivek bylo modelem chybně vyhodnoceno jako normální.

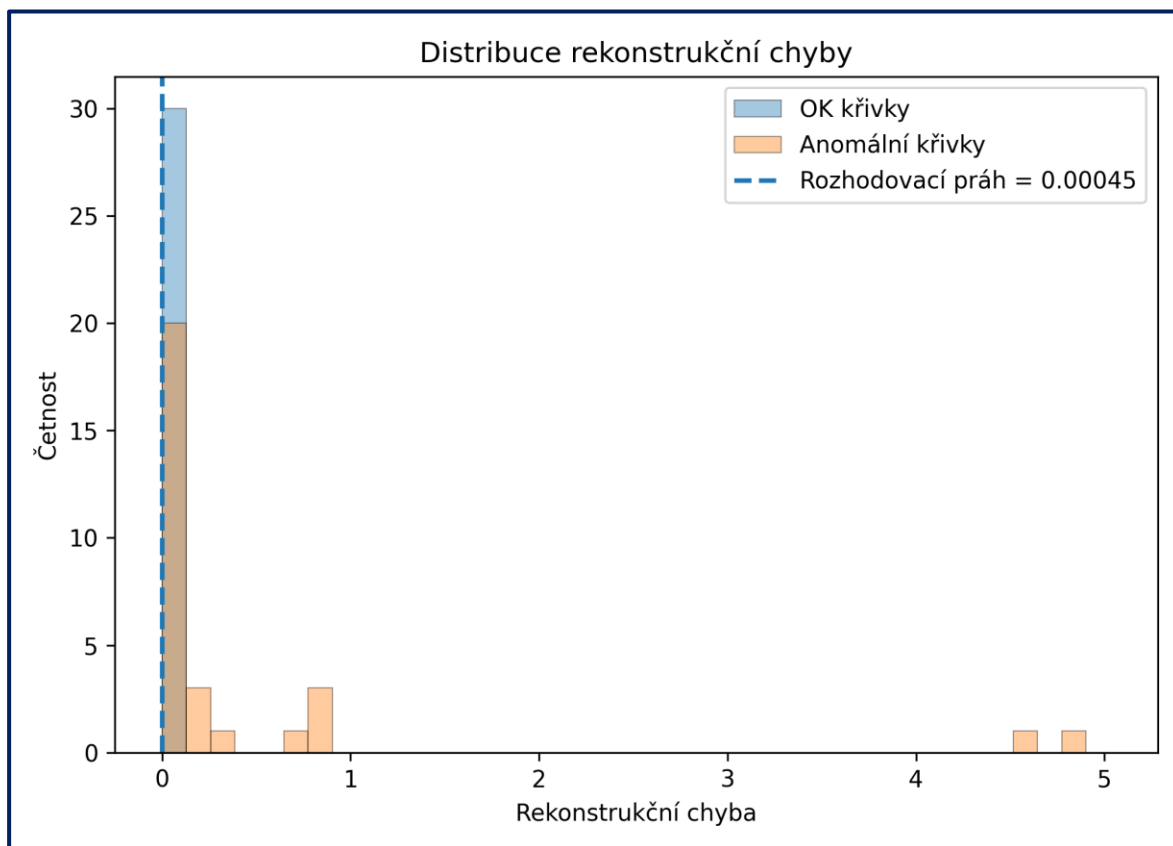
Tabulka 6.3: Metriky nasazeného modelu (zdroj: autorka, 2026)

	Metrika	Hodnota
	Rozhodovací práh (Threshold)	0,000446
	Přesnost klasifikace (Accuracy)	0,9167
	Přesnost pozitivní klasifikace (Precision)	1,0000
	Citlivost (Recall)	0,8333
	F1-skóre	0,9091
	Plocha pod ROC křivkou (AUC)	0,9522

Na základě dosažené hodnoty citlivosti modelu lze formulovat praktický předpoklad o jeho schopnosti zachytit anomální průběhy v reálném provozu. Model správně identifikoval 25 z 30 anomálních křivek v evaluační sadě, což odpovídá citlivosti 83,3 %. Za předpokladu, že rozložení anomálií v reálném provozu odpovídá rozložení v evaluační sadě, lze očekávat, že model zachytí přibližně čtyři z pěti anomálních průběhů.

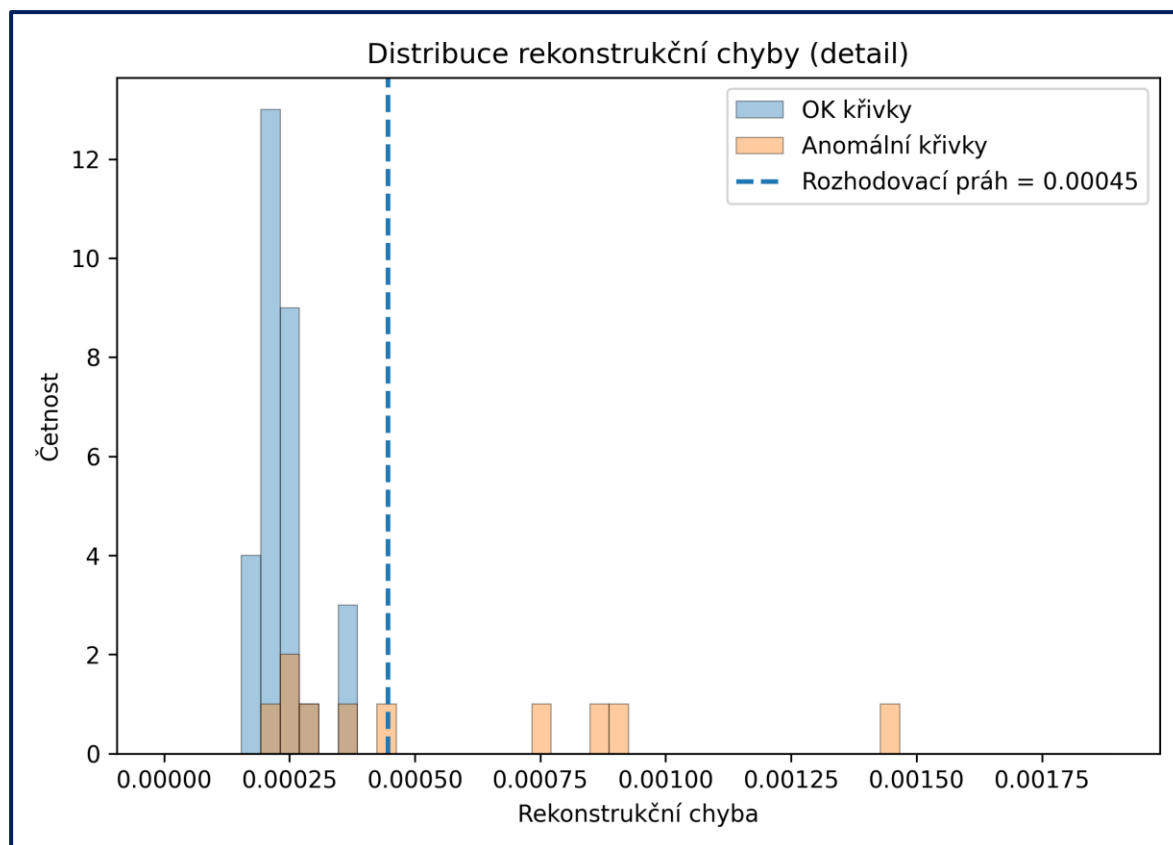
Model nevykazoval žádné falešně pozitivní případy – všechny normální křivky byly správně klasifikovány jako OK. Z pohledu výrobní praxe to znamená, že operátoři nejsou zatěžováni zbytečnými poplchy a každé upozornění modelu odpovídá skutečné odchylce.

Je nicméně nutné přistupovat k těmto odhadům **s obezřetností**, neboť **evaluační sada obsahovala pouze 30 NOK křivek**, přičemž každý jednotlivý zachycený nebo nezachycený kus představuje změnu citlivosti o přibližně 3,3 procentního bodu. Prezentované hodnoty jsou tedy orientační a pro robustnější odhad skutečné výkonnosti modelu v dlouhodobém provozu by bylo žádoucí ověření na větším a reprezentativnějším vzorku anomálních dat.



Obrázek 6.20: Distribuce rekonstrukční chyby pro normální a anomální momentové křivky (zdroj: autorka, 2026)

Obrázek 6.20 znázorňuje **rozdělení rekonstrukční chyby vypočtené pomocí konvolučního auto-encodeu pro dvě skupiny dat** – běžné a anomální křivky. Rekonstrukční chyba představuje rozdíl mezi původním průběhem momentové křivky a jejím průběhem rekonstruovaným modelem.

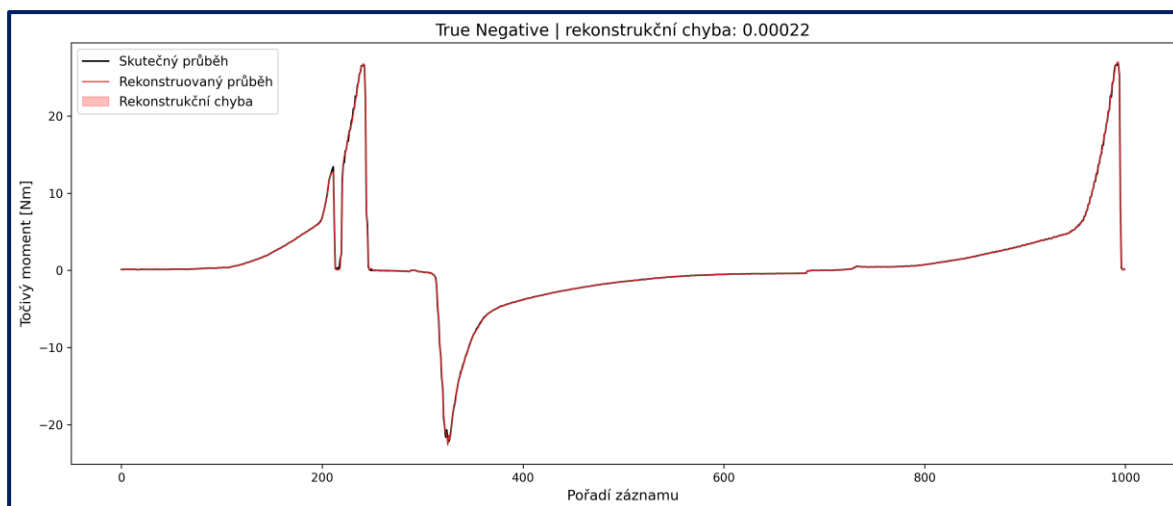


Obrázek 6.21: Detail distribuce rekonstrukční chyby pro normální a anomální momentové křivky v oblasti nízkých hodnot (zdroj: autorka, 2026)

V případě běžných křivek se hodnoty rekonstrukční chyby soustředí v úzké oblasti velmi nízkých hodnot, model tedy **dokáže tyto průběhy dobře rekonstruovat a jejich charakteristiky odpovídají vzorům naučeným během trénování**; naopak u anomálních křivek se rekonstrukční chyba zvyšuje a její hodnoty jsou rozprostřeny v širším rozsahu, v některých případech až o několik řádů vyšší než u běžných křivek. Je zřejmé, že kusy, které mají rekonstrukční chybu blízko nule, se od standardního průběhu vychylují jen málo a jsou to právě ty kusy, které mohou být velmi problematické, neboť je tradiční metody nejsou schopné označit za NOK.

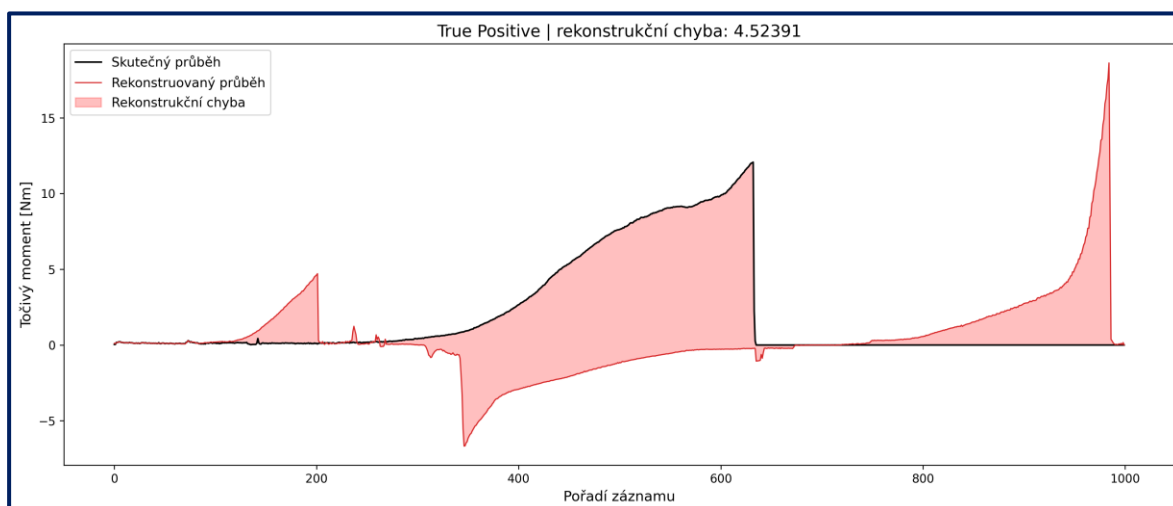
Křivky s rekonstrukční chybou nižší než práh stanovený z trénovacích dat jsou klasifikovány jako normální, zatímco křivky s vyšší chybou jsou označeny jako anomální. Z detailního grafu (viz **Obrázek 6.21**) je patrné, že mezi oběma skupinami existuje relativní separace, přesto však nedochází k jejich úplnému oddělení.

Dosavadním nedostatkem aplikace této metody v praxi může být **omezená reprezentativnost trénovacích dat**, omezený nebo ne zcela reprezentativní soubor evaluačních kusů či kombinace těchto dvou faktorů.



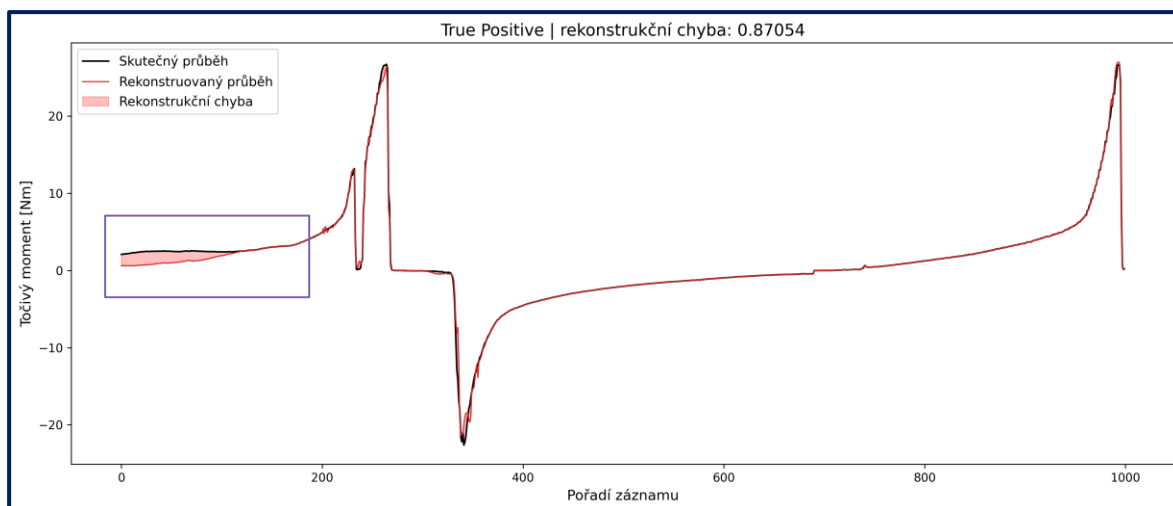
Obrázek 6.22: True negative křivka s velmi nízkou MSE (zdroj: autorka, 2026)

Obrázek 6.22 znázorňuje **reprezentativní příklad momentové křivky s normálním průběhem**, která byla, stejně jako všechny normální, správně modelem označena za neanomální, tedy true negative. Graf zachycuje průběh skutečného točivého momentu a jeho rekonstrukci vytvořenou konvolučním autoenkodérem. Na základě průběhu rekonstrukce lze říci, že autoenkodér kopíruje velmi dobře tvar původní křivky.



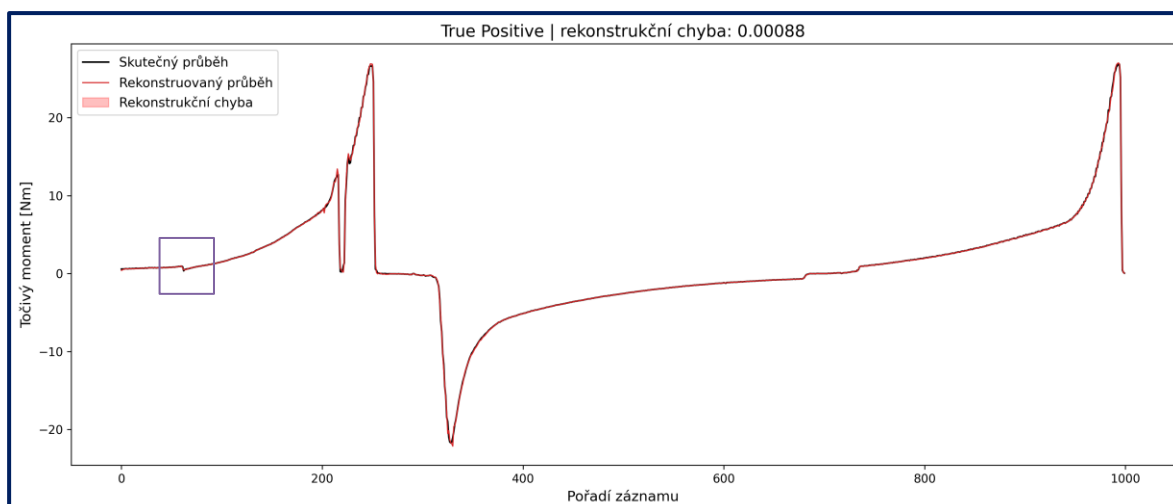
Obrázek 6.23: True positive křivka s vysokou MSE (zdroj: autorka, 2026)

Na **true positive křivce s velmi vysokou MSE** (viz Obrázek 6.23) lze jednoduše vidět červenou plochu, která **zobrazuje samotnou rekonstrukční chybu**, tedy kvantifikaci rozdílu mezi původním signálem a jeho aproximací generovanou modelem.



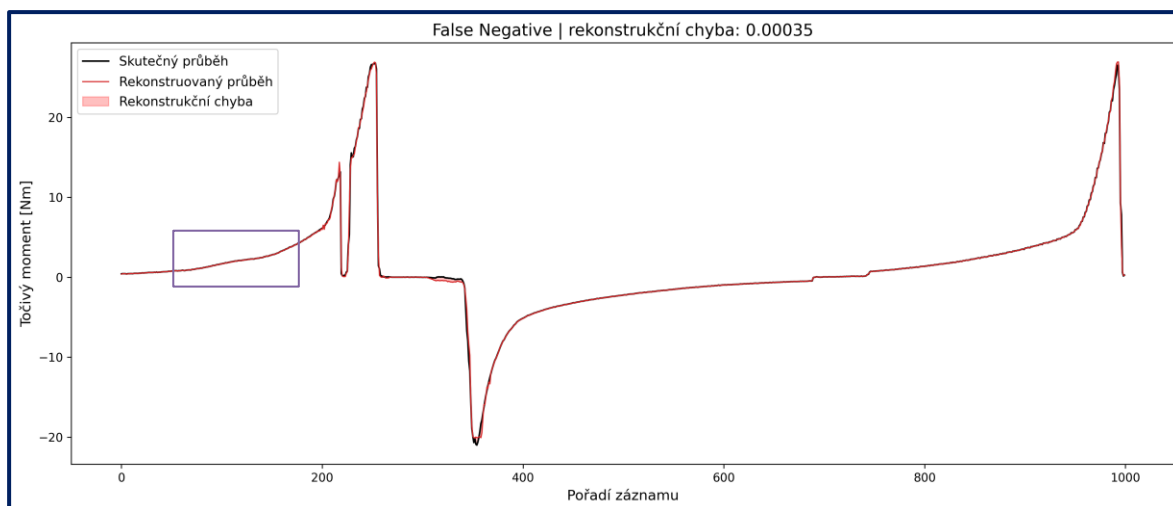
Obrázek 6.24: True positive křivka s relativně nízkou MSE (zdroj: autorka, 2026)

Na křivce **s relativně nízkou rekonstrukční chybou lze pozorovat typicky nestandardní průběh**, při kterém je zjevná chyba při zpracování konkrétního kusu – pravděpodobně by se ve žlutě zvýrazněné oblasti mělo jednat o skřípnutí těsnění, a tedy potenciální důvod k reklamaci.



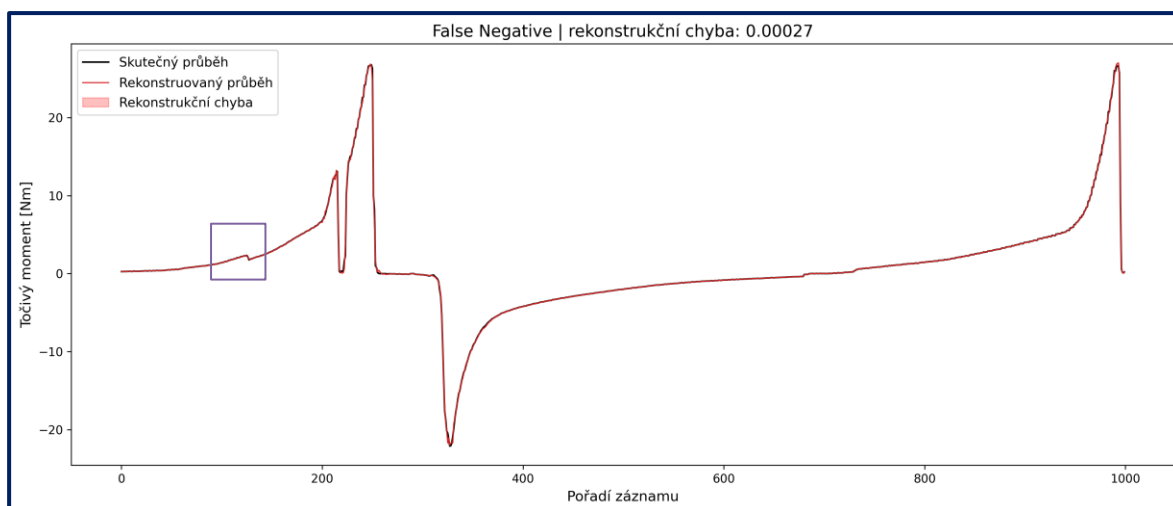
Obrázek 6.25: True positive křivka s MSE blízko k rozhodovacímu prahu (zdroj: autorka, 2026)

Výše vykreslený obrázek zobrazuje **anomální křivku, jež má pouze velmi malou, nezkušeným člověkem naprosto přehlédnutelnou odchylku**, která má nicméně velký význam z pohledu procesu – ne zcela plynulý nárůst na křivce mezi nultým a dvoustým bodem ve zvýrazněné oblasti naznačuje nějakou fyzickou nesrovnalost na dílu. Obdobné průběhy se staly předlohou pro tento výzkum, neboť dosud není zaveden systém, který by byl schopný na takovéto křivky se zdánlivě normálním průběhem upozorňovat.



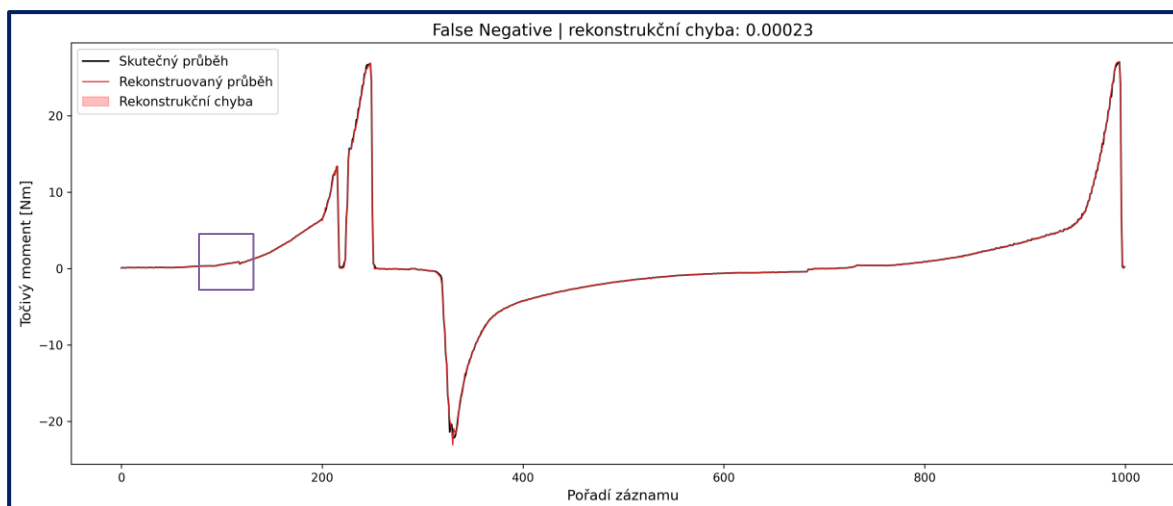
Obrázek 6.26: První false negative křivka (zdroj: autorka, 2026)

První křivka, která měla být označena za anomálii a nebyla, obsahuje **nestandardní vlnku mezi 50.–150. bodem. I** jako v předchozích a následujících případech jsou konkrétní anomálie zvýrazněny žlutým ohraničením.



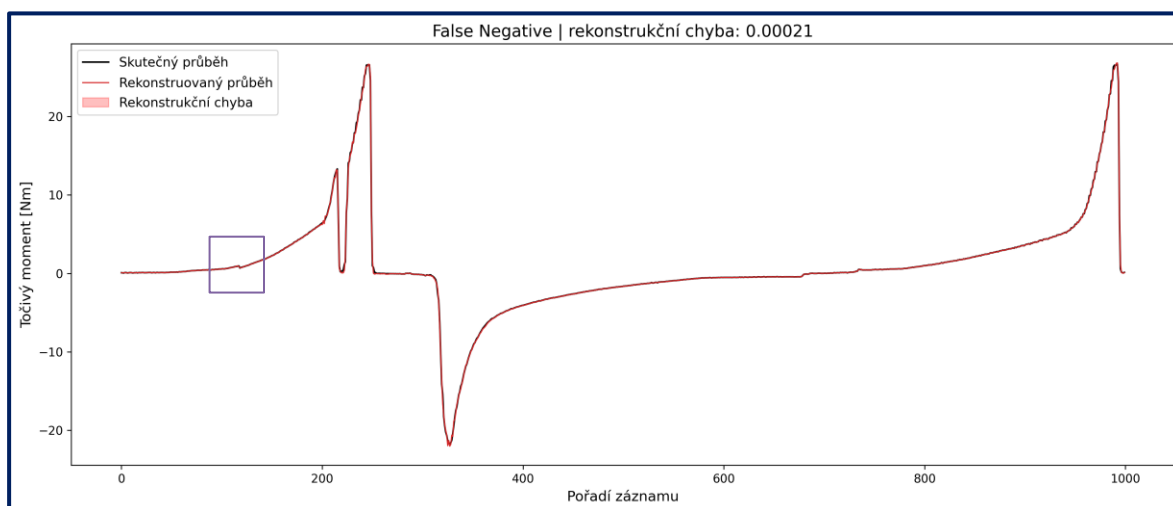
Obrázek 6.27: Druhá false negative křivka (zdroj: autorka, 2026)

Další obrázek s **false negative křivkou obsahuje mezi body 100–200 zub**, který značí anomální průběh šroubování kusu.



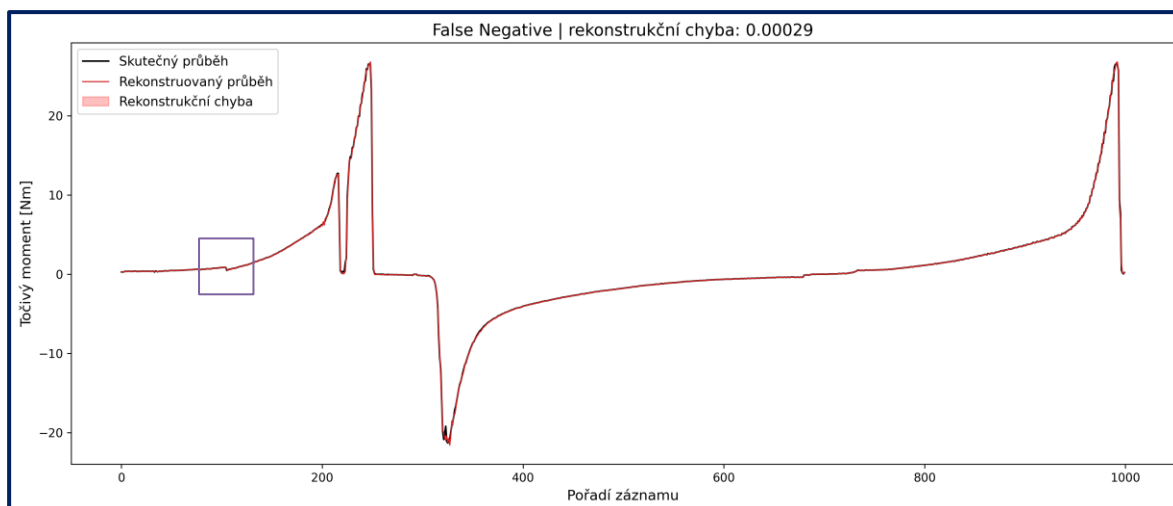
Obrázek 6.28: Třetí false negative křivka (zdroj: autorka, 2026)

Na třetí křivce označen expertem jako anomální lze pozorovat podivné chování okolo bodu 100.



Obrázek 6.29: Čtvrtá false negative křivka (zdroj: autorka, 2026)

Na čtvrté křivce je zobrazen velmi malý zub opět okolo bodu 100.



Obrázek 6.30: Pátá false negative křivka (zdroj: autorka, 2026)

Poslední false negative křivka ukazuje, stejně jako křivka předchozí, malý zub okolo bodu 100.

6.5 Závěry



- Aplikace navržených metod na křivková data ukazuje, že jejich **přínos** nespočívá pouze v možnosti detailněji analyzovat průběh procesu, ale také v tom, že **pomáhá řešit dva praktické problémy, které jsou ve výrobním prostředí úzce provázané** – vysoký objem ukládaných dat a současně obtížné rozpoznávání jemných, ale technologicky podstatných odchylek.
- **Ve fázi předzpracování** se ukazuje, že **část** původně ukládaných **informací nemá pro další analýzu skutečný význam**, a jejich vypuštění tak nevede ke ztrátě hodnoty, ale naopak ke zjednodušení a zpřehlednění datové struktury.
- **Následná komprese** potvrdila, že i **výrazná redukce** počtu bodů může při vhodně zvoleném nastavení **zachovat tvar křivky** v míře dostačující pro archivaci i pozdější zpětné vyhodnocení.
- Analýza významnosti jednotlivých úseků zároveň potvrdila expertní názor, dle kterého **informační hodnota není v rámci celé křivky rozložena rovnoměrně**, ale koncentruje se do určitých fází procesu.
- V oblasti **detekce anomálií** se pak ukázalo, že **model založený na učení běžného průběhu** je schopen zachytit i takové odchylky, které by při tradičním způsobu hodnocení zůstaly bez povšimnutí.
- **Křivková data** se tak v tomto pojetí neukazují pouze jako objemný záznam výrobního procesu, ale jako **zdroj informací, se kterým lze pracovat cíleněji** jak z hlediska ukládání, tak z hlediska včasné identifikace nestandardního chování.

7. Vyhodnocení přínosů



Účelem je detailně charakterizovat přínosy aplikovaných metod a řešení na skalárních i křivkových datech.

Vyhodnocení přínosů bylo provedeno ve třech rovinách odpovídajících jednotlivým navrženým optimalizačním opatřením. V případech, kde to charakter dostupných dat a metodického postupu umožňoval, byl **přínos kvantifikován a vyjádřen jako potenciální ekonomická úspora**. V případech, při nichž nebylo možné přínos jednoznačně vyčíslit, bylo hodnocení provedeno kvalitativně se zaměřením na praktickou využitelnost navrženého přístupu a jeho přínos pro práci s výrobními daty.

7.1 Detekce anomálií na skalárních datech

Navržený přístup **detekce anomálií vycházející ze skalárních výrobních dat** má z praktického hlediska především **koncepční a metodický význam**; vzhledem k absenci jednoznačně identifikovaných reklamovaných kusů však nebylo možné jeho výkonnost posoudit pomocí standardních klasifikačních metrik – hodnocení přínosu se proto opírá především o **schopnost metody zachytit odchylující se chování a nabídnout alternativní pohled na průběh výrobního procesu**. Metoda tak slouží primárně jako nástroj prioritizace; seřazuje výrobní kusy podle míry jejich odchylky od typického chování a usnadňuje tak výběr kandidátů pro expertní kontrolu.

Významnou předností tohoto řešení je jeho **nízká výpočetní náročnost**, neboť výpočet robustního z-skóre spolu s následnou agregací napříč parametry i stanicemi představuje operace, které rostou lineárně s objemem dat a nevyžadují žádné náročné trénování ani optimalizaci modelu – metoda je tak dobře škálovatelná a lze ji uvažovat i pro nasazení v prostředí s omezenými výpočetními zdroji, případně přímo pro on-line zpracování dat. Zároveň pracuje s relativně malým objemem vstupních dat, protože vychází pouze ze skalárních veličin, které jsou ve výrobní praxi běžně dostupné.

Podstatným přínosem je rovněž **schopnost převést heterogenní parametry na společnou škálu a následně hodnotit jejich odchylky v souvislostech**, a nejedná se tedy pouze o identifikaci extrémních hodnot jednotlivých veličin, ale také o zachycení situací, kdy se menší odchylky kumulují napříč více parametry či výrobními stanicemi – právě takové případy přitom často odpovídají reálnému chování procesu, kde se problémy neprojevují izolovaně, ale jako kombinace více dílčích vlivů.

Výsledný přehled s agregovaným anomaly score v současné fázi umožňuje **efektivní prioritizaci výrobních kusů pro další analýzu; v kontextu velmi nízkého podílu reklamací** – pohybujícího se přibližně na úrovni 0,01 % produkce – je totiž žádoucí zaměřit pozornost pouze na omezenou množinu potenciálně problematických případů. Metoda tak může sloužit jako nástroj předvýběru, který výrazně zúží objem dat určených k detailní kontrole a umožní cílenější práci s podezřelými kusy.

Je i tak nutné zdůraznit, že navržený přístup představuje spíše **první krok směrem k systematickému využití skalárních dat pro detekci anomálií**; jeho další rozvoj bude vyžadovat rozsáhlejší výzkum – zejména v oblasti validace na větším množství dat, nastavení rozhodovacích prahů a ověření vazby mezi vypočteným anomaly score a skutečnými výrobními defekty. Omezením zůstává také závislost na expertním výběru parametrů a stanovení vah jednotlivých stanic – jako jeden z prioritních kroků při dalším rozvoji metody je doporučena citlivostní analýza vah jednotlivých stanic, která by umožnila kvantifikovat vliv tohoto expertního rozhodnutí na výsledné pořadí kusů.

Navzdory těmto omezením lze hlavní přínos metody spatřovat v tom, že **otevívá možnosti detekce anomálií i v situacích, kde nejsou k dispozici rozsáhlé datové sady**, detailní označení vad nebo dostatečná výpočetní kapacita pro nasazení komplexních modelů – ukazuje tak, že i s využitím relativně jednoduchých statistických nástrojů je možné získat smysluplný pohled na odchylky ve výrobním procesu a vytvořit základ pro další, pokročilejší analytické přístupy.

7.2 Komprese křivkových dat

Pro odhad ekonomického přínosu komprese křivkových dat byl nejprve určen objem generovaných dat. Průměrná velikost jedné křivky z pilotní linky se pohybuje okolo 75 kB a na jedné výrobní lince vzniká ročně přibližně půl milionu těchto záznamů, což odpovídá zhruba 37,5 GB dat za rok. Při

předpokládané době uchování 10 let tak dosahuje celkový objem uložených křivkových dat přibližně 375 GB na jednu linku.

Při praktickém ověření se ukázalo, že **navržený způsob komprese při zachování potřebných meta-dat snižuje objem dat v průměru o 81 %**. U pilotní linky, kde roční objem dat dosahuje přibližně 37,5 GB, to znamená úsporu zhruba 30,4 GB ročně; při ceně 0,72 € za 1 GB odpovídá tato hodnota přibližně 21,9 € za rok. Na úrovni celého pracoviště jsou však objemy dat nesrovnatelně vyšší – denně vzniká přibližně jeden milion křivek o průměrné velikosti okolo 1 MB, tedy zhruba 1 TB dat denně, což představuje přibližně 365 TB za rok. Při zachování stejné míry komprese by tak bylo možné ročně snížit objem dat přibližně o 296 TB, což při uvedené ceně odpovídá úspoře přibližně 213 tisíc € ročně jen na křivkových datech.

Vedle samotné redukce objemu dat je však přínos navrženého řešení širší; komprese neznamená pouze úsporu kapacity úložiště, zároveň **snižuje nároky na přenos dat mezi jednotlivými systémy a zjednodušuje jejich další zpracování**. Výhodou přístupu je rovněž skutečnost, že **redukce probíhá řízeným způsobem, tedy se zachováním klíčových charakteristik signálu a potřebných meta-dat**; nedochází tak k pouhému „odstranění dat“, ale k jejich efektivnějšímu uložení při zachování informační hodnoty.

Uvedený výpočet vychází výhradně z křivkových dat, která z hlediska datového objemu nepředstavují nejnáročnější kategorii výrobních záznamů. Na řadě linek jsou paralelně ukládána i obrazová data a videozáznamy, jejichž velikost je o řád vyšší; **ověření kompresního přístupu na momentových křivkách tak slouží spíše jako doklad toho, že systematická redukce datového objemu je realizovatelná** při zachování podstatné informační hodnoty signálu. U objemnějších datových typů by se stejný princip promítl do výrazně vyšších úspor.

7.3 Detekce anomálií na křivkových datech

Následující výpočet představuje ilustrativní odhad za předpokladu ideálních podmínek. Skutečný přínos **závisí na charakteru konkrétních reklamací, reprezentativnosti evaluační sady a stabilitě výrobního procesu v čase**. Pro robustní vyčíslení by bylo nutné dlouhodobé nasazení modelu v reálném provozu.

Na základě dostupných provozních údajů lze orientačně vyčíslit potenciální ekonomický přínos nasazení modelu pro detekci anomálií na momentových křivkách – průměrný náklad na jednu reklamaci činí minimálně 25 000 € a ročně je u pilotní linky č. 2 evidováno přibližně 17 reklamačních případů, celkový roční náklad je tedy okolo 425 000 €.

Za předpokladu, že model dosahuje v reálném provozu citlivosti odpovídající výsledkům na evaluační sadě, tedy přibližně 83,3 %, bylo by možné zachytit zhruba 14 z 17 reklamačních případů ročně ještě před odesláním výrobků zákazníkovi – **potenciální roční úspora by tak dosahovala přibližně 354 000 €**. Zbývající přibližně tři případy by i nadále mohly projít nezachyceny, **s odhadovaným nákladem okolo 71 000 €**. Při statistické nejistotě dané malým rozsahem evaluační sady je nutné zmínit, že se skutečná hodnota zachycených reklamací může výrazně lišit.

Vedle samotného finančního efektu má navržený přístup **největší potenciál z pohledu organizace kontroly; model lze využít jako nástroj pro včasnou identifikaci problémů**, takže namísto zpětného řešení reklamací je možné podezřelé kusy zachytit už během výroby, případně bezprostředně po jejím dokončení – tím se nejen snižuje riziko odeslání vadného výrobku zákazníkovi, ale zároveň vzniká prostor pro rychlejší reakci na vznikající odchylky v procesu.

Výhodou tohoto přístupu je také jeho **relativní obecnost, neboť metoda není pevně svázána s konkrétním parametrem ani typem výrobku**, ale vychází z tvaru signálu jako takového – při dostatečně reprezentativních datech ji lze aplikovat i na jiné typy křivek nebo na další části výrobního procesu. Současně jde o řešení, které nevyžaduje rozsáhlé množství označených dat, protože staví na učení z běžného chování – to je v průmyslových podmínkách zásadní, jelikož vadné kusy tvoří jen velmi malý podíl produkce.

Je však potřeba brát v úvahu, že uvedené **vyčíslení přínosu je pouze orientační** a vychází z několika zjednodušujících předpokladů, zejména že každá reklamace odpovídá anomálii zachytitelné modelem a že výkonnost modelu zůstane v dlouhodobém provozu na úrovni dosažené při evaluaci. Ve skutečnosti může být přínos ovlivněn řadou dalších faktorů – například charakterem konkrétních reklamací, kvalitou a reprezentativností trénovacích dat nebo přirozenou variabilitou výrobního procesu v čase; právě tyto oblasti by proto měly být dále ověřeny při případném nasazení v praxi.

7.4 Závěry



- Z uvedených výsledků je patrné, že **přínos navržených přístupů se projevuje v různých rovinách**; zatímco komprese dat přináší přímé úspory v oblasti ukládání a práce s daty, detekce anomálií má hlavní hodnotu ve včasné identifikaci nestandardního chování a podpoře rozhodování ve výrobě.
- Lze vidět, že **nejde pouze o finanční efekt, ale i o zefektivnění práce s daty a cílenější kontrolu**; zároveň je však nutné výsledky chápat jako orientační, jejich plné ověření by vyžadovalo dlouhodobé nasazení v reálném provozu.
- Uvedený postup se byl **prakticky ověřen** a ukazuje se, že je **reálně použitelný**. Použité metody dávají smysluplné výsledky a potvrzují, že zvolený směr optimalizace není jen teoretický, ale lze ho aplikovat i na skutečná výrobní data.
- Zároveň je ale potřeba počítat s tím, že **šlo o ověření na omezeném vzorku**. Pokud by se měl postup nasadit v praxi, bylo by nutné ho dál testovat na delším časovém období a na širším spektru dat; teprve potom by se dalo spolehlivěji říct, jak stabilně funguje v běžném provozu a jaký má skutečný přínos.

8. Závěr

Cílem bylo navrhnout a **ověřit přístupy k efektivnějšímu zpracování výrobních dat s důrazem na snížení jejich objemu a současně na využití jejich informační hodnoty** pro detekci nestandardního chování procesu. V rámci řešení byly analyzovány jak skalární, tak křivkové datové struktury, přičemž navržené postupy se zaměřily na jejich praktickou využitelnost v podmínkách reálné výroby.

V oblasti skalárních dat byl navržen přístup založený **na robustním vyhodnocení odchylek jednotlivých parametrů a jejich následné agregaci** napříč výrobními stanicemi. Výsledky ukázaly, že takto definované anomaly score umožňuje identifikovat kusy, u nichž dochází ke kumulaci odchylek v průběhu výroby, a poskytuje tak nový pohled na chování procesu. Přínos této metody spočívá především v její jednoduchosti a nízké výpočetní náročnosti.

V případě křivkových dat byly řešeny dvě oblasti – jejich objem a možnost detekce anomálií. Navržený způsob komprese prokázal, že je možné významně snížit datový objem při zachování podstatné informační hodnoty signálu; zároveň byl implementován i model pro detekci anomálií založený na rekonstrukční chybě, který umožňuje zachytit nestandardní průběhy momentových křivek a potenciálně tak předcházet vzniku reklamací. Orientační vyčíslení přínosu naznačuje, že i částečné zlepšení včasné detekce může vést k významným ekonomickým úsporám.

Zdroje

- Ali, A., Laghari, A. A., Kandhro, I. A., Kumar, K., & Younus, S. (2024). Systematic analysis of on-premise and cloud services. *International Journal of Cloud Computing*. https://www.researchgate.net/publication/379970310_Systematic_Analysis_of_On_Premise_and_Cloud_Services
- Amazon Web Services. (2024). *Siemens Energy Builds Industrial IoT Platform and Drives Smart Manufacturing Using AWS IoT*. AWS Customer Stories. <https://aws.amazon.com/solutions/case-studies/siemens-energy-video-case-study/>
- Badman, A., & Kosinski, M. (2024). *What is big data?* <https://www.ibm.com/think/topics/big-data>
- Dignan, L. (2025). *A look at Walmart's platform approach to data, AI, optimization*. Constellation Research. <https://www.constellationr.com/insights/news/look-walmarts-platform-approach-data-ai-optimization>
- Douglas, D. H., & Peucker, T. K. (1973). ALGORITHMS FOR THE REDUCTION OF THE NUMBER OF POINTS REQUIRED TO REPRESENT A DIGITIZED LINE OR ITS CARICATURE. *Cartographica*, 10(2), 112–122. <https://doi.org/10.3138/FM57-6770-U75U-7727>
- Facebook. (2011). *Facebook Launches Open Compute Project*. Meta Newsroom. <https://about.fb.com/news/2011/04/facebook-launches-open-compute-project/>
- Goodfellow, I., Bengio, Y., & Courville, A. (2016a). Autoencoders. In *Deep Learning*. MIT Press. <https://www.deeplearningbook.org/contents/autoencoders.html>
- Goodfellow, I., Bengio, Y., & Courville, A. (2016b). Convolutional Networks. In *Deep Learning*. MIT Press. <https://www.deeplearningbook.org/contents/convnets.html>
- Kamalov, F., Sulieman, H., Alzaatreh, A., Emarly, M., Chamlal, H., & Safaraliev, M. (2025). Mathematical Methods in Feature Selection: A Review. *Mathematics*. <https://doi.org/10.3390/math13060996>
- Mehra, V., & Hasegawa, R. (2023). Supporting power grids with demand response at Google data centers. *Google Cloud Blog*. <https://cloud.google.com/blog/products/infrastructure/using-demand-response-to-reduce-data-center-power-consumption>
- Meiners, M., Mayr, A., & Franke, J. (2020). Process curve analysis with machine learning on the example of screw fastening and press-in processes. *Procedia CIRP*, 97, 166–171. <https://doi.org/10.1016/j.procir.2020.05.220>
- Naeem, M., Jamal, T., Diaz_Martinez, J., Butt, S. A., Montesano, N., Tariq, M. I., De-la-Hoz Franco, E., & De-La-Hoz-Valdiris, E. (2022). Trends and Future Perspective Challenges in Big Data. In *Advances in Intelligent Data Analysis and Applications*. https://doi.org/10.1007/978-981-16-5036-9_30
- O'Donovan, P., Leahy, K., Bruton, K., & T. J. O'Sullivan, D. (2015). *Big data in manufacturing: A systematic mapping study*. https://www.researchgate.net/publication/281646704_Big_data_in_manufacturing_a_systematic_mapping_study
- Olszewski, R. T. (2001). *Generalized Feature Extraction for Structural Pattern Recognition in Time-Series Data* [Disertační práce, Carnegie Mellon University, School of Computer Science]. https://www.researchgate.net/publication/228558473_Generalized_feature_extraction_for_structural_pattern_recognition_in_time-series_data
- Quantum News. (2025). *Deepmind AI Cuts Google Data Center Cooling Bill By 40%, Revolutionizing Energy Efficiency*. Quantum Zeitgeist. <https://quantumzeitgeist.com/deepmind-ai-cuts-google-data-center-cooling-bill-by-40-revolutionizing-energy-efficiency/>
- Rousseeuw, P. J., & Hubert, M. (2018). Anomaly detection by robust statistics. *WIREs Data Mining and Knowledge Discovery*. <https://doi.org/10.1002/widm.1236>
- SNECI. (2025). *Automotive industry: Definition, players, challenges and advice*. <https://www.sneci.com/en/industries/automobile/>
- Strohbach, M., Daubert, J., Ravkin, H., & Lischka, M. (2016). Big Data Storage. In *New Horizons for a Data-Driven Economy: A Roadmap for Usage and Exploitation of Big Data in Europe*. Springer Open. <https://doi.org/10.1007/978-3-319-21569-3>

Tan, Y., Hahn, O., & Du, F. (2005). Process Monitoring Method with Window Technique for Clinch Joining. *ISIJ International*, 45(5). https://www.jstage.jst.go.jp/article/isijinternational/45/5/45_5_723/_pdf

TensorFlow. (2024). *Intro to Autoencoders | TensorFlow Core*. TensorFlow.org. <https://www.tensorflow.org/tutorials/generative/autoencoder>

Yeo, I.-K., & Johnson, R. A. (2000). A new family of power transformations to improve normality or symmetry. *Biometrika*, 87(4), 954–959. <https://doi.org/10.1093/biomet/87.4.954>