



Scaling AI adoption in insurance by increasing trust

Presented by Shameek Kundu
and David Marock

February 2022



Speaker bios



David Marock – CEO, NED, Chair, Senior advisor & Angel investor

David was, most recently, Group CEO of Charles Taylor plc, a global insurance services & technology business; he joined from specialist insurer, Beazley plc, where he was Group COO and on the Beazley Furlonge Ltd board. Prior to that, he was an Associate Principal at McKinsey.

David is now a senior advisor to McKinsey, along with various PE firms & tech-related start-ups; finally, he is an angel investor in tech firms. He is also a NED at Standard life Savings, one of the UK's largest investment platforms and PremFina, a fintech transforming the premium finance market. He has been a NED at Standard Life Assurance, one of the UK's largest insurers; Fadata, an international insurance technology company; and The Standard Club, a global P&I insurer. David is a Fellow of the Faculty and Institute of Actuaries.



Shameek Kundu - Head of Financial Services and Chief Strategy Officer, TruEra

Shameek is Chief Strategy Officer and Head of Financial Services at TruEra, a software firm dedicated to improving the quality of AI models. He has spent most of his career driving responsible adoption of data analytics (including AI) in the financial services industry

Shameek sits on the [Bank of England's AI Public-Private Forum](#), and was part of the Monetary Authority of Singapore's Steering Committee on [Fairness, Ethics, Accountability and Transparency in AI](#). Most recently, Shameek was Group Chief Data Officer at Standard Chartered Bank, where he helped the bank explore and adopt AI in multiple areas (e.g., credit, financial crime compliance, customer analytics).

Who are we

- **Early stage software company** with presence in California, Seattle, London, New York and Singapore
- Solely dedicated to **AI/ Machine Learning Quality** (explainability, stability, reliability, fairness and other aspects)
- Experienced, **hands-on leadership team** from a mix of Financial Services (SCB, Visa), Big Tech (Google, Microsoft), Academia (Carnegie Mellon) and startups
- Multiple live banking and insurance clients, including one **public Global systemically important bank** (SCB) and **a top North American insurer**
- Extensive engagement with financial services industry



Monetary Authority
of Singapore



BANK OF ENGLAND



INSTITUTE OF
INTERNATIONAL
FINANCE

THE FINTECH TIMES

FINANCIAL TIMES

NU PROPERTY
CASUALTY360°

insuranceday



INSURANCE
THOUGHT
LEADERSHIP.COM

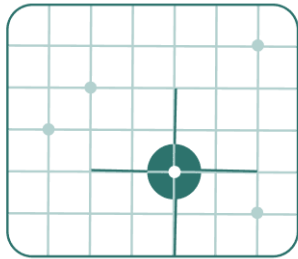
InsuranceERM
The online resource for enterprise risk management

Scaling AI adoption in insurance by increasing trust

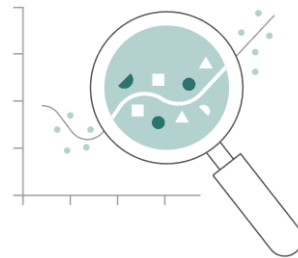
Topics of discussion:

- Ways in which AI can transform the entire insurance value chain
- Current status of AI adoption across the insurance market
- Best practice approaches to capture AI's full potential

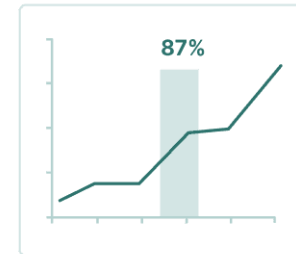
AI holds great promise for insurers



Supplement traditional models for underwriting and pricing of risk



Automating aspects of claims and fraud management



Improve customer service and automate Back-Office Ops

Two major changes turbo-charging the current AI opportunity for insurers



Data explosion

Explosion in data available to support decision-making

- Unstructured documents – e.g., detailed information filed to regulators by publicly listed companies – easy to analyse
- Structured data –e.g., satellite imagery - more widely accessible
- New categories of data – e.g., personal health data from wearable devices – now available



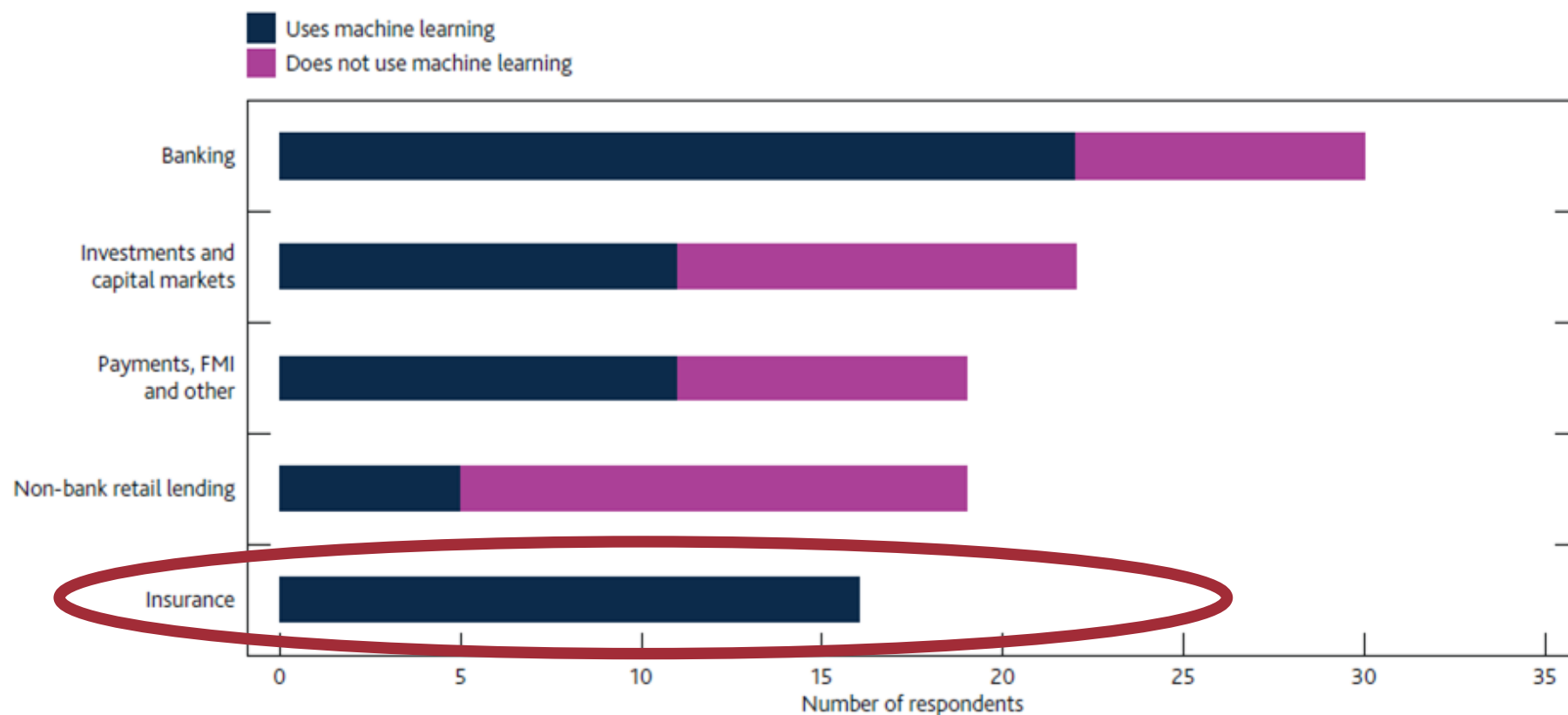
Maturity of AI techniques

Dramatic increase in maturity in the last few years

- Computing power and data storage has become substantially cheaper
- AI modeling tools now both more accessible and powerful

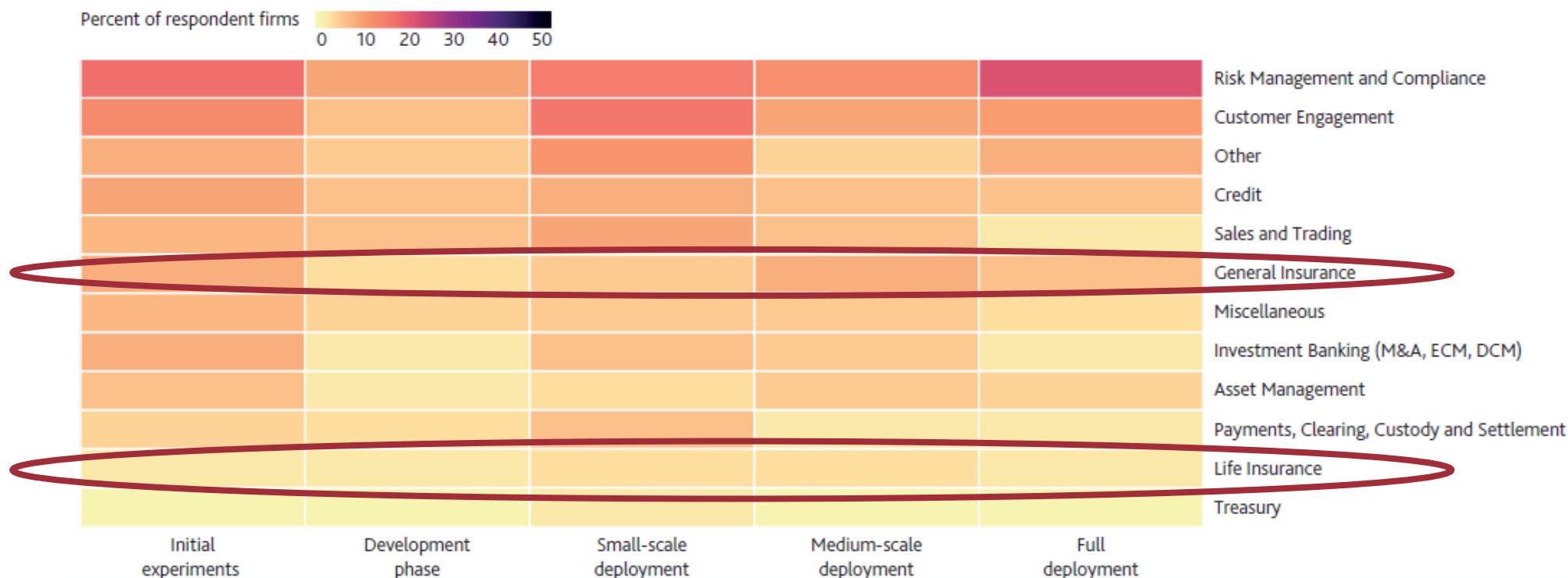
AI adoption in financial services

While insurers appear active in adopting AI/ ML ...



AI adoption in financial services

...adoption is still in its early stages across much of insurance



Different types of insurers are at different stages of adoption , for example:

- Personal line carriers, e.g., auto and health insurers, at the forefront of early adoption
- More specialized segments, e.g., large corporate risk and specialty insurers, have been experimenting with AI

Six things that make it difficult to scale AI



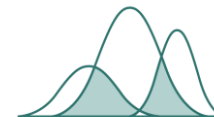
Explainability

What factors are driving the behaviours of the automated claims assessment system?



Fairness/ Unjust Bias

Why are women getting quoted lower premiums than men?



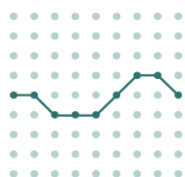
Stability

Are historical pricing models ageing well with increased frequency of climate change related events?



Conceptual Soundness

Is the model consistent with our domain knowledge?



Reliability

How confident is the model in its predictions?



Data Bias

Does the training data accurately reflect the population?

Explainability: ML Models are often black boxes

Training Data & Labels

“Convertible”

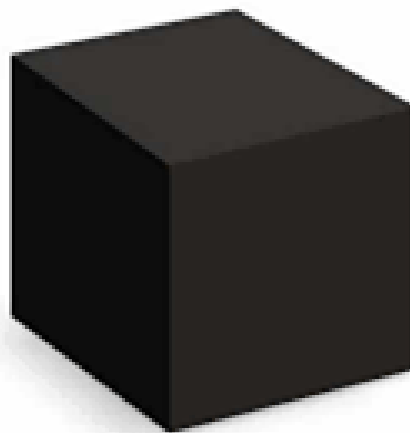


“Sports Car”



Model is a black box

Pattern Example: “Roof” pixels correlated with “sports car”



Hard to understand results

New unlabeled data

“?”



“?”



“Convertible”



“Sports Car”



Accurate explanations are critical to establish trust with model stakeholders.

Stability: ML Models are at greater risk from data drift

Training Data & Labels

“Convertible”

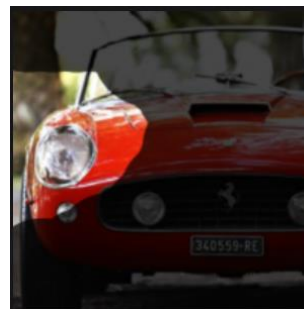


“Sports Car”



Models can learn patterns that become outdated over time

“Round headlights” correlated with “convertible”



≡ “Convertible”

Data changes/drifts surface model instability & reduce performance

New unlabeled data

“?”



“Sports Car”



“?”



“Sports Car”



i

In the covid era, enterprises need to understand and react to model drift.

Overfitting: ML Models have greater risk of overfitting to training Data

Training Data & Labels

“Convertible”



“Sports Car”



Model “memorizes” patterns in specific training instances

“Model learns certain license plates associated with convertibles”



≡ “Convertible”

Memorization doesn’t generalize well reducing model performance

New unlabeled data

“?”



“Sports Car”



“?”



“Sports Car”



i

Detecting and addressing overfitting is a key step during development.

Fairness: ML Models are at greater risk of reinforcing biases

Training Data & Labels
“Cooking”



COOKING	
ROLE	VALUE
Agent	Woman
Food	Fruit
Heat	0
Tool	Knife
Place	Kitchen



COOKING	
ROLE	VALUE
Agent	Woman
Food	Meat
Heat	Stove
Tool	Spatula
Place	Outside

Models can learn patterns that reflect historical biases

“Woman” = “Cooking”

Model predictions are biased
New unlabeled data



COOKING	
ROLE	VALUE
Agent	Woman
Food	0
Heat	Stove
Tool	Spatula
Place	Kitchen



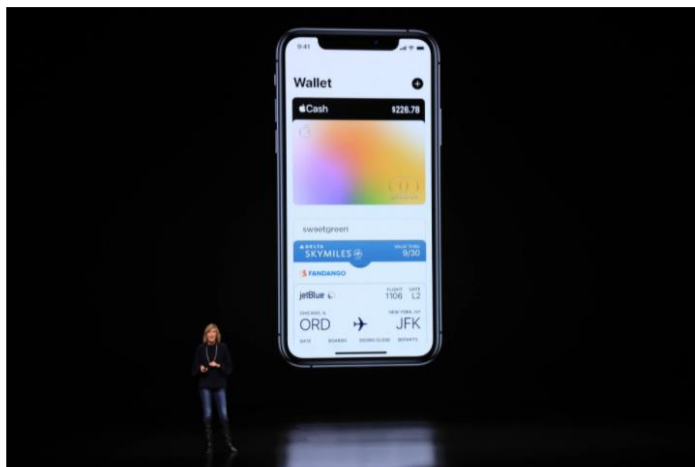
The use of alternative data could sometimes help mitigate historical bias.

Reputational impact from poorly designed/explained AI

The New York Times

Apple Card Investigated After Gender Discrimination Complaints

A prominent software developer said on Twitter that the credit card was “sexist” against women applying for credit.

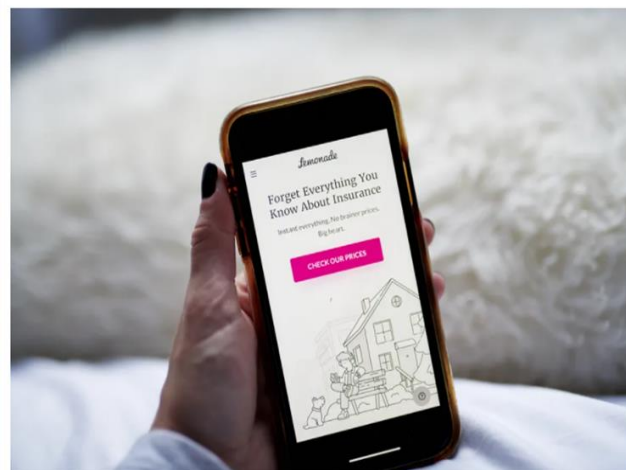


Jennifer Bailey, vice president of Apple Pay. Regulators are investigating Apple Card's algorithm, which is used to determine applicants' creditworthiness. Jim Wilson/The New York Times

A disturbing, viral Twitter thread reveals how AI-powered insurance can go wrong

Lemonade tweeted about what it means to be an AI-first insurance company. It left a sour taste in many customers' mouths.

By Sara Morrison | May 27, 2021, 1:30pm EDT



Lemonade wants you to forget everything you know about insurance. | Gabby Jones/Bloomberg via Getty Images

Lemonade, the fast-growing, machine learning-powered insurance app, put out a real lemon of a Twitter thread on Monday with a proud declaration that its AI analyzes videos of customers when determining if their claims are fraudulent. The company has been trying to explain itself and its business model — and fend off serious accusations of bias, discrimination, and general creepiness — ever since.

Amazon's sexist AI recruiting tool: how did it go so wrong?



Julien Lauret

Follow

Aug 15, 2019 · 18 min read ★



Photo by Tim Gouw on Unsplash

Last year, Reuters broke the news that Amazon had been working on a secret AI recruiting tool that showed bias against women. I found it interesting as a case study of an AI project with broad implications for business people and machine learning professionals. After all, everyone has

Quality of AI key topic for global financial regulators

Insurance regulators have recognised risks and set expectations for insurers to implement AI responsibly:

- Principles on AI governance published by the US National Association of Insurance Commissioners (NAICS) in 2020
- European Insurance and Occupational Pensions Authority (EIOPA) in 2021.

“Providers of high-risk AI systems shall...have a quality management system in place...”

“AI use should be fair, ethical, accountable and transparent”

“...risks related to data, models, firm-level accountability and governance and entire financial system”

“...data that present greater consumer protection risks warrant more robust compliance management... [including] appropriate testing, monitoring and controls...”



“...tools in transparency and explainability ...suggestions on how the governance of the use of AI should be organised to safeguard sound use of AI”

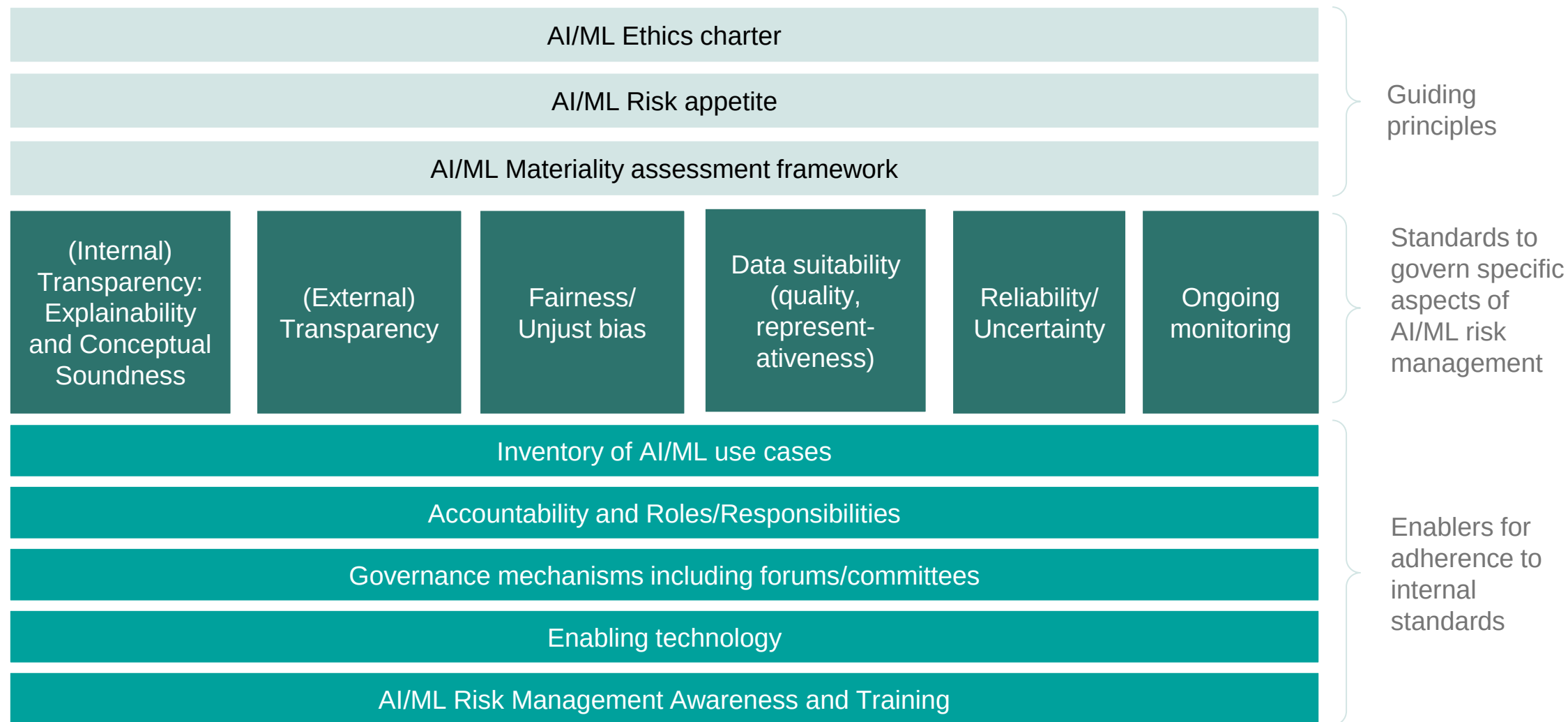
“...a systematic risk management approach to each phase of the AI system life cycle on a continuous basis to address risks related to AI systems”



So how can insurers overcome these obstacles and realize the full value of AI?

- The risks related to AI should not become a reason to stall greater adoption
- Progressive insurers have ensured this by:
 - a. Establishing internal policies/ standards/ guidelines for responsible use of AI
 - b. Increasing internal awareness and capability across business, technology and data science
 - c. Introducing tools to improve AI quality; e.g.,
 - i. Explain the underlying drivers of the model outputs accurately
 - ii. Identify potential instances of unjust bias and recommend mitigating steps
 - iii. Monitor and troubleshoot models' performance on an ongoing basis

Foundational risk management framework for AI in insurance



How new technology enables insurers to trust AI

Targeted Marketing

- Assess whether marketing approach is promoting products unsuitable to customers' circumstances and/ or needs
- Determine whether certain segments, however defined, are being actively, possibly inappropriately, targeted or avoided

Claims management

- Monitor the predictive power and stability of the claims triage model on an ongoing basis
- Help investigators understand why particular claims have been flagged as fraudulent

Underwriting

- Understand what drives underwriting models' decisions as to which market segments to cover, along with which exclusions or limits to apply
- Ensure that insurance underwriting model decisions can be explained easily to regulators

Investment process

- Understand drivers of investment advice models, thereby identifying potential biases
- Monitor the reliability of investment advice over time

Technical pricing

- Confirm rating models are not using - either directly or as a proxy - data points that are not regulatory compliant
- Obtain regulatory approval for rating models by clearly explaining how premiums are being calculated by the model

Operational automation

- Provide early warnings when data drift is likely to impact the accuracy of models used to automate back-office processes
- Diagnose quickly root causes for performance degradation in operational efficiency models

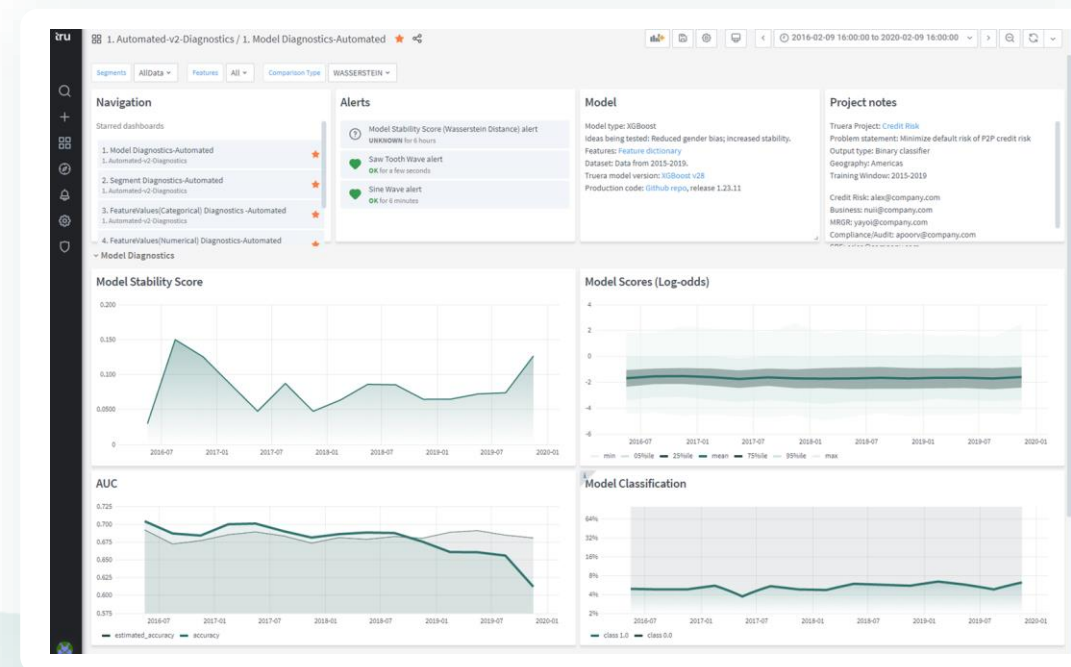
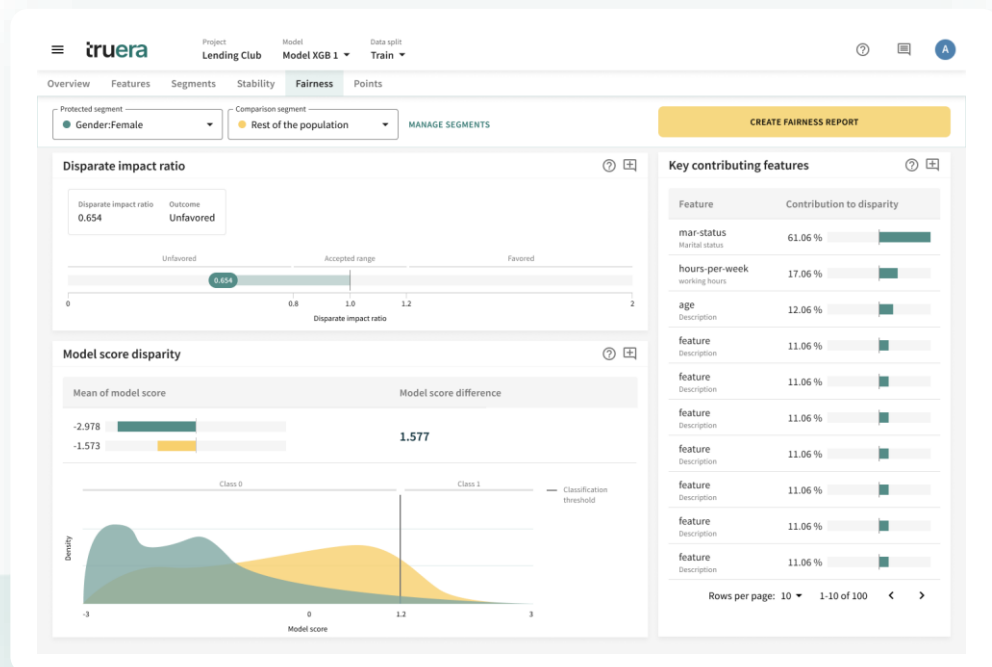
TruEra's AI solutions address AI trust issues – upfront and on an ongoing basis

TruEra Diagnostics

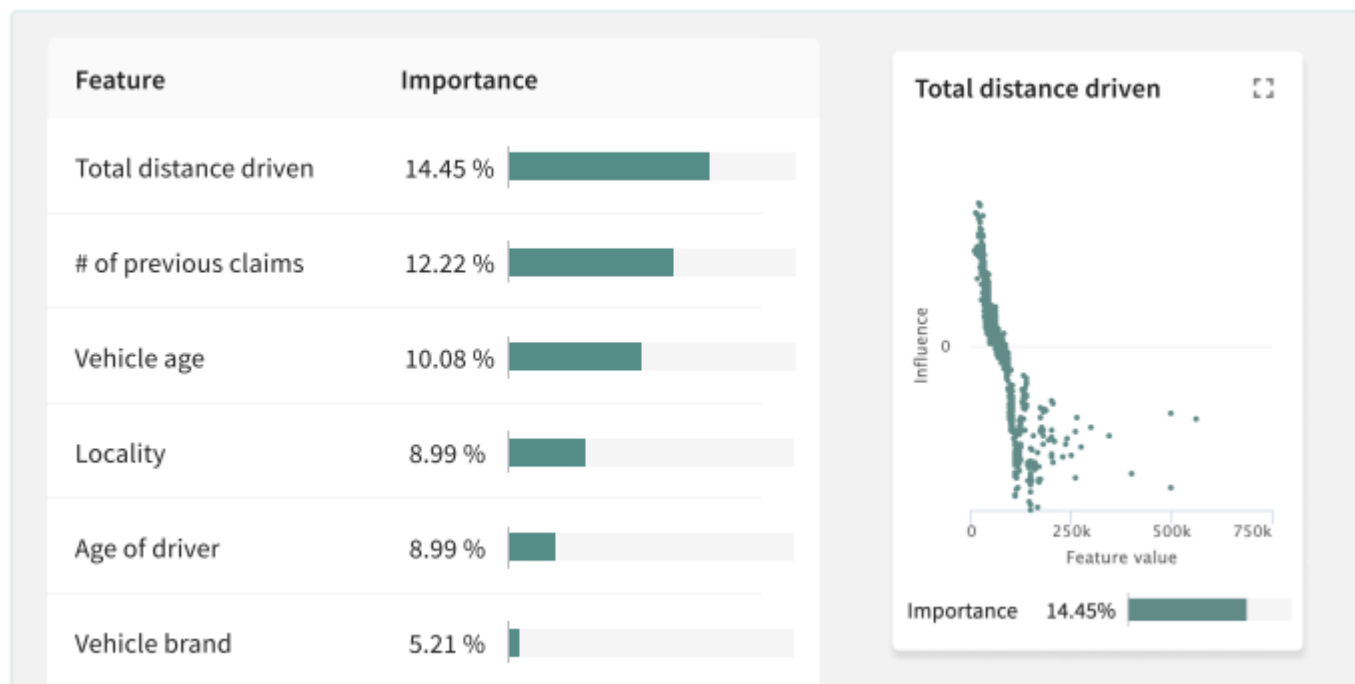
Evaluate, explain and gain adoption during development

TruEra Monitoring

Supervise & debug models post deployment



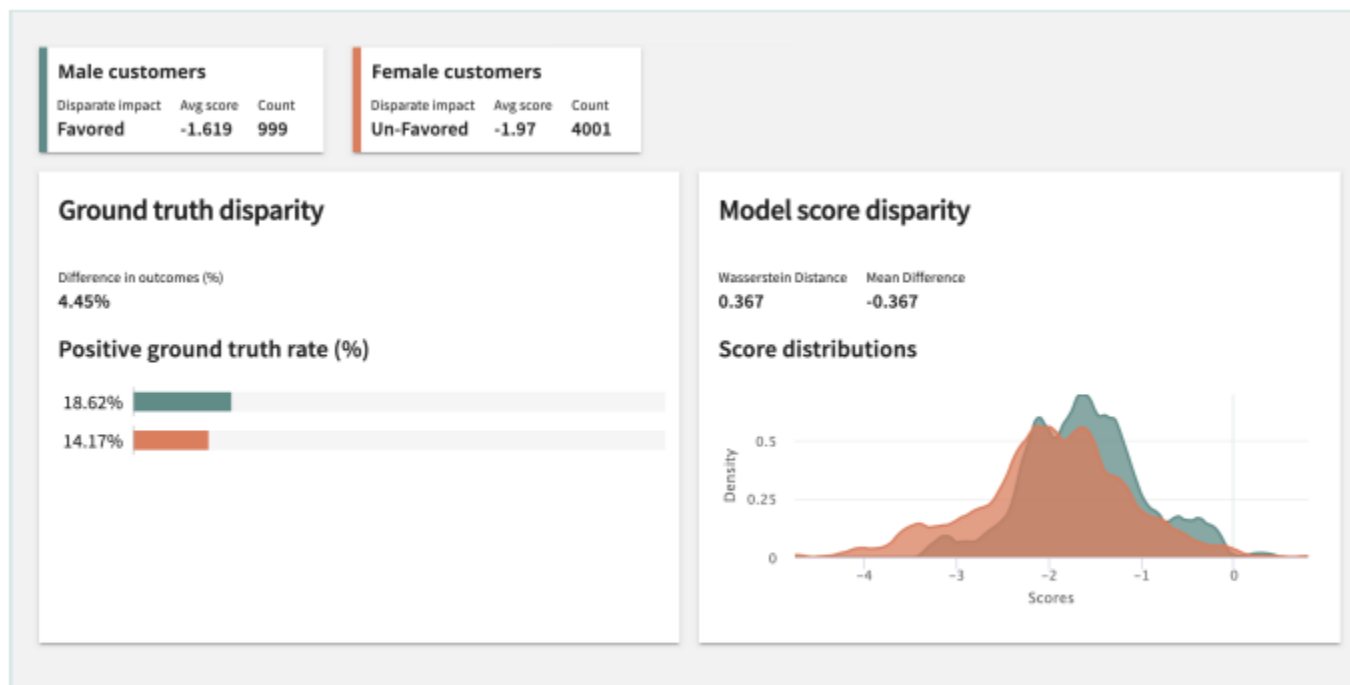
Example: TruEra breaks open the AI/ML “blackbox”



Trust & Explainability

- More accurate & faster explanations (10-10,000 times faster than SHAP)
- Feature grouping / reason code support APIs
- Conceptual & segment analysis to build trust

Example: TruEra enables detection and treatment of unjust bias in AI/ML models

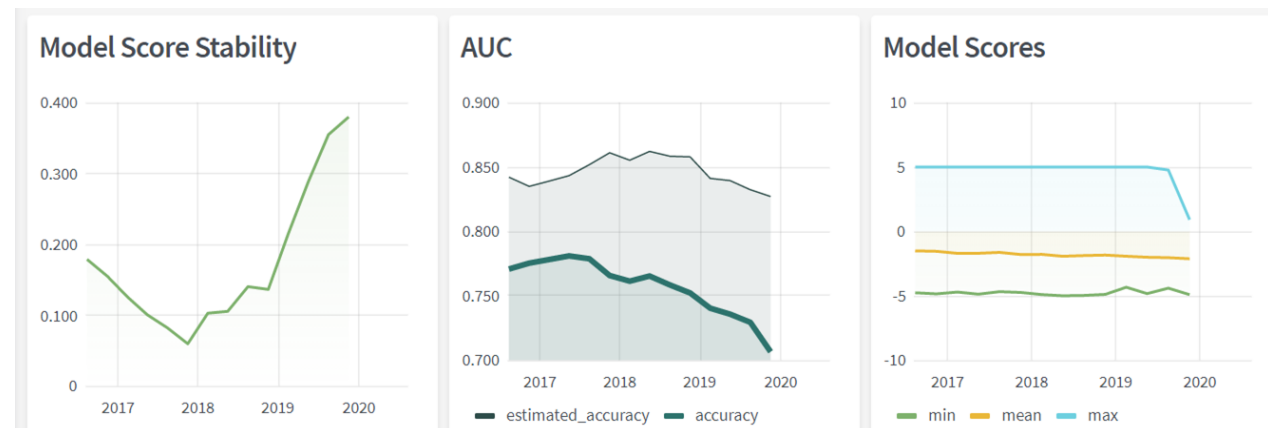


Assessment of fairness/ bias

- Comprehensive capabilities to enable assessment of ML models against regulatory or internal standards

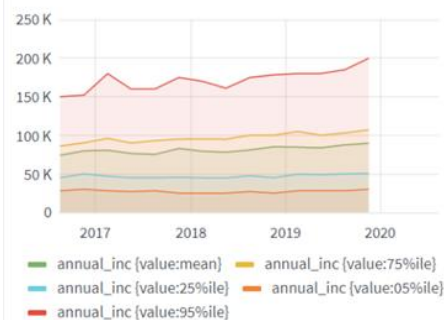
Fairness - Differences in model's treatment of different groups (e.g., gender)

Example: TruEra Monitoring



▼ Numeric Feature Values: annual_inc

Feature Values: annual_inc



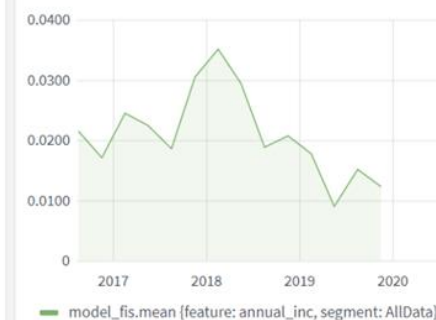
Feature Values(Low): annual_inc



Feature Values(High): annual_inc



Feature Influence Stability: ann...



Scaling AI adoption in insurance by increasing trust

Summary of discussion:

- Various material ways in which AI can transform entire insurance value chain
- AI adoption across the insurance market is still relatively early stage
- Best practice approaches to capture AI's full potential involve both processes & tech tools