

Fake Lepton Estimation

F9 Strategic Meeting Krvavec • September 28th, 2026



Jan Gavranovič

Jožef Stefan Institute and University of Ljubljana



Why fake leptons, and why data-driven

- Any analysis that selects leptons (e , μ , τ_{had}) also selects objects that only *look* like prompt leptons: **fake (misidentified or non-prompt) leptons**.
- In many final states (multilepton, same-sign, low lepton p_T , ...) they are a **dominant background**.
- **MC cannot model fakes reliably** (tails of jet fragmentation and detector response, enormous MC statistics needed) → **data-driven estimation is mandatory**.
- Example in this talk: **electrons**, faked by sources with **different fake rates**:
 - light-flavour jets misidentified as electrons,
 - photon conversions ($\gamma \rightarrow e^+ e^-$),
 - non-prompt electrons from semileptonic heavy-flavour (b , c) decays.
- ATLAS uses a family of data-driven techniques (fake factors, matrix method, template fits, ABCD, ...) → choice depends on the analysis.

This talk: fake-factor (FF) method for electrons → F as a density ratio → ML estimator.

The fake-factor method

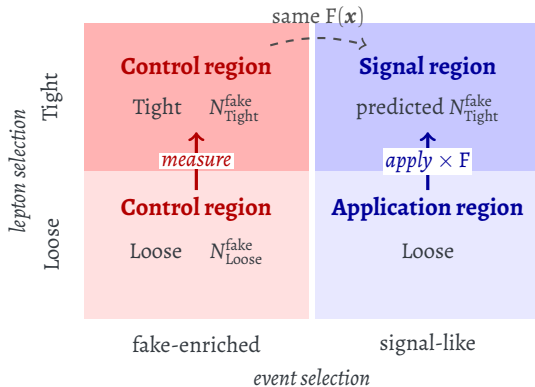
- Split by **lepton selection**: **Tight** (analysis ID + isolation) and **Loose** (relaxed, fail Tight).
- Measured** in a fake-enriched **control region (CR)**, real leptons subtracted with MC:

$$F = \frac{N_{\text{Tight}}^{\text{data}} - N_{\text{Tight}}^{\text{real}}}{N_{\text{Loose}}^{\text{data}} - N_{\text{Loose}}^{\text{real}}} \Big|_{\text{CR}}.$$

- Applied** to the Loose events of the SR (Tight SR data never used):

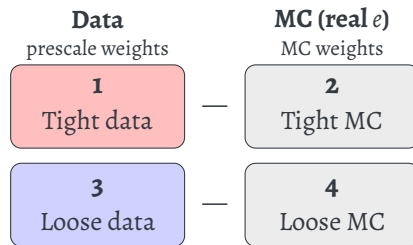
$$N_{\text{Tight}}^{\text{fake}} = F \times (N_{\text{Loose}}^{\text{data}} - N_{\text{Loose}}^{\text{real}}) \Big|_{\text{SR}}.$$

- Assumption: **same F in CR and SR at fixed $\mathbf{x}$** $\rightarrow$ parametrized in $\mathbf{x}$ (e.g. $p_{\text{T}} \times |\eta| \times E_{\text{T}}^{\text{miss}}$).



Example: electron fakes in a fake-enriched region

- Events collected with **prescaled electron triggers**.
- Price: every data event carries its **prescale as a weight**, from 1 up to $10^5 \rightarrow$ **weighted** densities.
- **b -jet veto** suppresses $t\bar{t}$ and heavy-flavour decays.
- **Real electrons** (e.g. $W(e\nu)+\text{jets}$, $Z(ee)+\text{jets}$) in both subregions are subtracted with MC $\rightarrow$ **four samples**.
- Baseline: **binned** F in $p_T \times |\eta| \times E_T^{\text{miss}}$ $\rightarrow$ reference for the ML estimate.

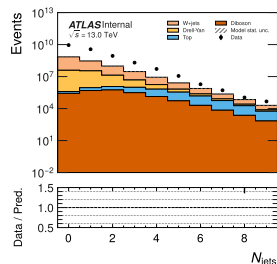
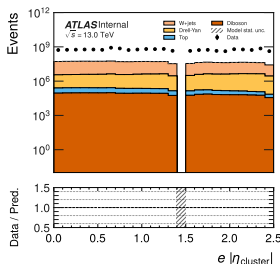
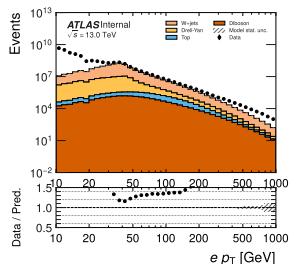


$$F = \frac{N_1 - N_2}{N_3 - N_4} \Big|_{\text{CR}}$$

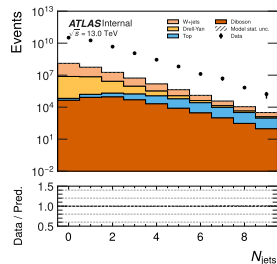
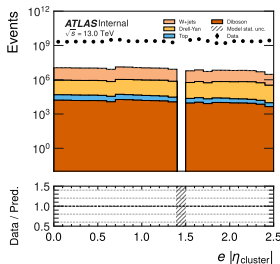
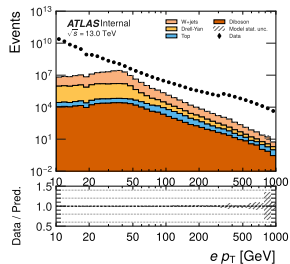
Dataset	Trigger	electron p_T	electron $ \eta $	$N_{b\text{-jet}}$	F binning
Full Run 2	prescaled electron	≥ 10 GeV	< 2.47 , excl. 1.37–1.52	$= 0$	$p_T \times \eta \times E_T^{\text{miss}}$

Tight and loose distributions

Tight (numerator)



Loose (denominator)



From binned to continuous fake factors

- The fake factor is a **ratio of real-subtracted densities** in the *Tight* and *Loose* regions:

$$F = \frac{N_f^T}{N_f^L} = \frac{N_{\text{data}}^T - N_{\text{MC}}^T}{N_{\text{data}}^L - N_{\text{MC}}^L} \longrightarrow F(\mathbf{x}) = \frac{N_f^T(\mathbf{x})}{N_f^L(\mathbf{x})}.$$

- Traditionally binned in a few variables (for electrons p_T , $|\eta|$, E_T^{miss}): **binning choice** (bias vs. variance), statistics per bin, empty or **negative bins** after the subtraction, jumps at bin edges, and **at most a few dimensions**.
- **Idea: estimate the density ratio with neural networks trained directly on data** (“data-based inference”) [JHEP 04 (2026) 188]: a **continuous, unbinned, per-event** fake factor in a high-dimensional feature space.
- Event weights (MC modeling, prescales) and the real-electron subtraction are handled **natively in the training** $\rightarrow$ **one network, one training**.

The building blocks

- **Count density** $n(\mathbf{x})$: expected events per unit volume of $\mathbf{x}$.
- **Weighted density** $D(\mathbf{x})$: the same as $n(\mathbf{x})$ with weights.
- **Factorization**: count density times conditional mean weight:

$$D(\mathbf{x}) = n(\mathbf{x}) \bar{w}(\mathbf{x}) , \quad \bar{w}(\mathbf{x}) = \mathbb{E}[w \mid \mathbf{x}] \quad (\text{in one bin: } \sum_i w_i = N\bar{w}) .$$

- **Classifier ratio**: train a classifier on sample a ($y = 1$) vs. sample b ($y = 0$) with binary cross-entropy and *no weights*, its logit q_{ab} gives:

$$\ell(q, y) = -y \log \sigma(q) - (1 - y) \log (1 - \sigma(q)) \quad \Longrightarrow \quad e^{q_{ab}(\mathbf{x})} = \frac{n_a(\mathbf{x})}{n_b(\mathbf{x})} .$$

- **Mean weight regression**: least squares on the events of sample a ,

$$\ell(\bar{w}, w) = (\bar{w} - w)^2 \quad \Longrightarrow \quad \bar{w}_a(\mathbf{x}) = \mathbb{E}[w \mid \mathbf{x}, a] ,$$

negative weights only enter inside a square $\rightarrow$ safe.

- **No weights in n , the sign lives in $\bar{w} \rightarrow$ signed weights never enter a classifier loss.**

Decomposition of the fake factor

- The goal is to express F using only ratios, via the identity $u - v = u(1 - v/u)$.
- Divide numerator and denominator by their own weighted data density:

$$F(\mathbf{x}) = \frac{D_1(\mathbf{x}) - D_2(\mathbf{x})}{D_3(\mathbf{x}) - D_4(\mathbf{x})} = \frac{D_1(\mathbf{x}) (1 - D_2(\mathbf{x})/D_1(\mathbf{x}))}{D_3(\mathbf{x}) (1 - D_4(\mathbf{x})/D_3(\mathbf{x}))} = r_{13} \frac{1 - 1/r_{12}}{1 - 1/r_{34}},$$

where every ratio has the same form,

$$r_{ab}(\mathbf{x}) = \frac{D_a(\mathbf{x})}{D_b(\mathbf{x})}, \quad \text{so} \quad r_{13} = \frac{D_1}{D_3} \quad \text{and} \quad r_{12} = \frac{D_1}{D_2} \quad \text{and} \quad r_{34} = \frac{D_3}{D_4}.$$

- We can now factorize each weighted density using:
 - $D(\mathbf{x}) = n(\mathbf{x})\bar{w}(\mathbf{x})$ to get $D_a = n_a\bar{w}_a$,
 - $e^{q_{ab}}(\mathbf{x}) = n_a(\mathbf{x})/n_b(\mathbf{x})$ to get the three classifier heads $e^{q_{12}} = n_1/n_2$, $e^{q_{34}} = n_3/n_4$ and $e^{q_{13}} = n_1/n_3$,
- *The final result is a decomposition into seven parts:*

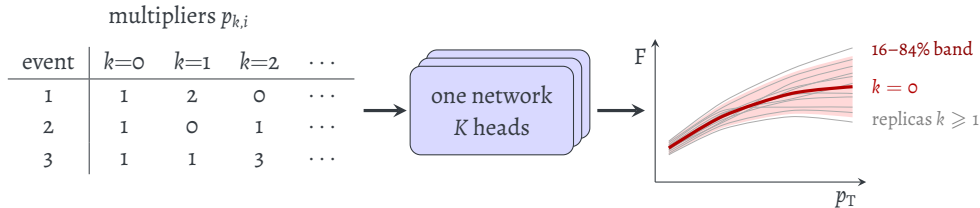
$$r_{ab}(\mathbf{x}) = e^{q_{ab}}(\mathbf{x}) \frac{\bar{w}_a(\mathbf{x})}{\bar{w}_b(\mathbf{x})} \quad \text{for} \quad (a, b) \in \{(1, 2), (3, 4), (1, 3)\}.$$

Poisson bootstrap for uncertainties

- Statistical uncertainty on $F(\mathbf{x})$ from a **Poisson bootstrap vectorized inside the model**.
- Each event i gets a multiplier per channel k : $p_{0,i} = 1$ (nominal) and $p_{k,i} \sim \text{Poisson}(1)$ for $k \geq 1$, fixed by hashing the event number with k .
- The multipliers turn each channel into a **statistically independent pseudo-dataset**.
- Every loss term is a $p_{k,i}$ -weighted mean of a per-event quantity v_i over a sample S :

$$\langle v \rangle_k^S = \frac{\sum_{i \in S} p_{k,i} v_i}{\sum_{i \in S} p_{k,i}}.$$

- The loss averages the K channels $\rightarrow$ K **independent estimators** (K heads) trained in **one forward and backward pass** $\rightarrow$ band = 16th/84th percentiles over $k \geq 1$.



The loss, term by term

Seven heads per channel k : logits q_{12}, q_{34}, q_{13} and mean weights $\bar{w}_1, \dots, \bar{w}_4$.

- **Classifier**, one per ratio $(a, b) \in \{(1, 2), (3, 4), (1, 3)\}$: binary cross-entropy ℓ_{ab} on $S_{ab} = a \cup b$, **no event weights**:

$$C_{ab} = \langle \ell_{ab} \rangle_k^{S_{ab}} \implies e^{q_{ab}} = \frac{n_a}{n_b}.$$

- **Mean weight**, one per sample $a = 1, \dots, 4$: relative squared error:

$$R_a = \left\langle \left(\frac{\bar{w}_a - w}{\text{sg}[\bar{w}_a]} \right)^2 \right\rangle_k^a \implies \bar{w}_a = \mathbb{E}[w | \mathbf{x}, a].$$

- **Positivity**, one per subtraction $(1, 2), (3, 4)$: penalize $r_{ab} < 1$, i.e. more MC than data:

$$P_{ab} = \gamma \left\langle \max(0, -\log r_{ab})^2 \right\rangle_k^{S_{ab}} \implies r_{ab} \geq 1 \text{ (no negative fake yield)}.$$

The total loss

- **Total loss:** add the terms per channel, average over the K bootstrap channels,

$$L = \frac{1}{K} \sum_{k=0}^{K-1} \left[\underbrace{\sum_{(a,b)} C_{ab,k}}_{\text{ratios } n_a/n_b} + \underbrace{\sum_{a=1}^4 R_{a,k}}_{\text{mean weights}} + \underbrace{\sum_{(1,2), (3,4)} P_{ab,k}}_{\text{positive subtraction}} \right],$$

one backward pass trains all seven heads in all K channels together.

- **After training** the heads give the fake factor, channel by channel:

$$r_{ab} = e^{q_{ab}} \frac{\bar{w}_a}{\bar{w}_b} \quad \longrightarrow \quad F = r_{13} \frac{1 - 1/r_{12}}{1 - 1/r_{34}},$$

channel $k = 0$ is the nominal F , channels $k \geq 1$ give the band.

Applying F needs loose data only

- Applying F to the real-subtracted loose sample gives the **predicted tight-fake density**:

$$P(\mathbf{x}) = F(\mathbf{x}) (D_3(\mathbf{x}) - D_4(\mathbf{x})) .$$

- Factor the **loose subtraction** with the same identity as before, $u - v = u(1 - v/u)$:

$$D_3 - D_4 = D_3(1 - 1/r_{34}) .$$

- Insert the decomposition of F and **cancel the loose factor**:

$$P = r_{13} \frac{1 - 1/r_{12}}{1 - 1/r_{34}} \times D_3(1 - 1/r_{34}) = D_3 r_{13} (1 - 1/r_{12}) .$$

- The only density left is D_3 , so at event level this is the **loose data** sample weighted by $w_i r_{13}(\mathbf{x}_i) (1 - 1/r_{12}(\mathbf{x}_i))$: **no MC events are needed** (*data-driven*).

From three logs to F

- The **three logs** are put together from the heads:

$$v_{ab} = \log r_{ab} = q_{ab} + \log \bar{w}_a - \log \bar{w}_b .$$

- Each factor $1 - 1/r_{ab}$ of the decomposition is the **fake share**:

$$F = r_{13} \frac{1 - 1/r_{12}}{1 - 1/r_{34}} , \quad 1 - \frac{1}{r_{ab}} = 1 - e^{-v_{ab}} = -\text{expm1}(-v_{ab}) ,$$

so F is evaluated as

$$F = e^{v_{13}} \cdot \frac{-\text{expm1}(-v_{12})}{\max(-\text{expm1}(-v_{34}), \varepsilon)} .$$

- Applied to loose data: **the denominator cancels**, so nothing divides by the small, noisy subtraction and ε is never needed.

F vs p_T in slices of $|\eta|$ and E_T^{miss}

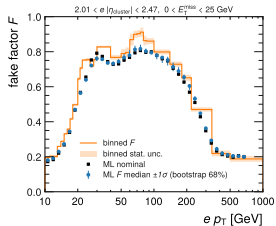
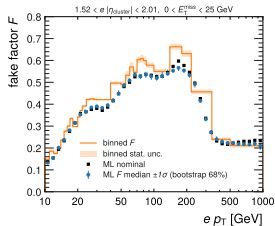
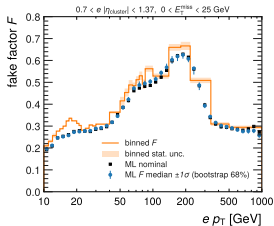
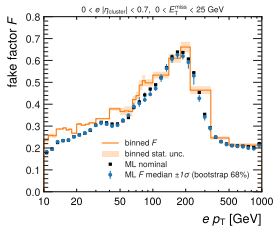
$0 < |\eta| < 0.7$

$0.7 < |\eta| < 1.37$

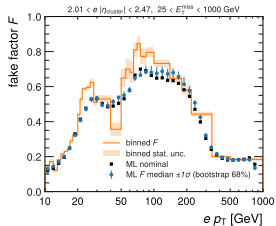
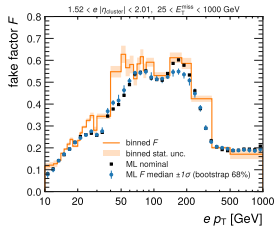
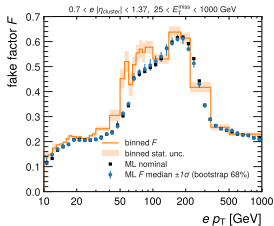
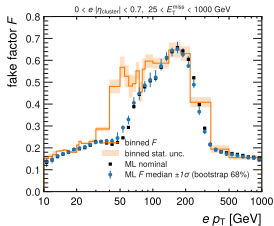
$1.52 < |\eta| < 2.01$

$2.01 < |\eta| < 2.47$

$E_T^{\text{miss}} < 25 \text{ GeV}$

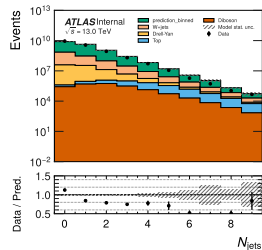
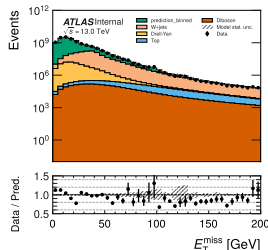
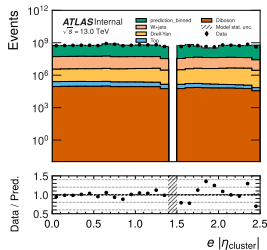
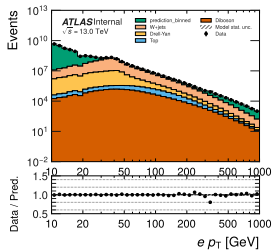


$E_T^{\text{miss}} > 25 \text{ GeV}$

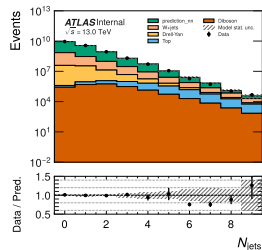
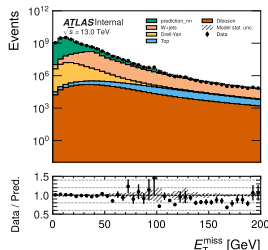
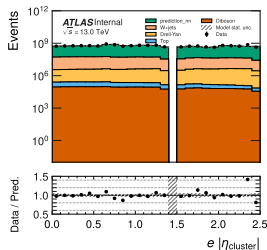
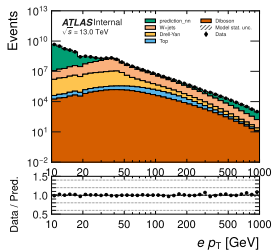


Closure: binned vs ML

Binned F



ML F



Conclusion

- Fake-factor method with a **continuous, unbinned, per-event** F , shown for electron fakes in a fake-enriched region.
- **One network, one training**: seven heads on a shared backbone give F directly, instead of three separately trained classifiers.
- Weighted densities factorize as $D = n \bar{w}$: classifiers learn count ratios, regressions learn mean weights $\rightarrow$ **no signed weight ever enters a classifier loss**.
- Uncertainties come for free from the vectorized **Poisson bootstrap**, in the same forward and backward pass.
- The closure is good in p_T , $|\eta|$, E_T^{miss} and N_{jets} , and matches the binned fake factor.
- **Code**: gitlab.cern.ch/ijs-f9-ljubljana/nnsf.
- **Questions and feedback are very welcome.**