

Frequently asked questions

Can you recommend some reading material on the background of the NEGF method?

For a first reading on NEGF, there are several good introductions published by [Supriyo Datta](#). Other books or articles on NEGF usually require advanced quantum mechanics or quantum field theory.

- The book *Quantum Transport: Atom to Transistor* by S. Datta is an introduction to NEGF which is not too advanced.
- There is a more advanced book by S. Datta, *Electronic Transport In Mesoscopic Systems*. Chapter 8 is a good introduction.
- Another rather simple article is [Nanoscale device modeling: the Green's function method](#), S. Datta, *Superlattices and Microstructures* **28**, 253 (2000)
- There is a review article on QCLs where T. Kubis describes in Section VIII the NEGF method. [Modeling techniques for quantum cascade lasers](#), C. Jirauschek, T. Kubis, *Applied Physics Reviews* **1**, 011307 (2014)

How much memory do I need?

We recommend 16 GB RAM. 8 GB is sufficient for some input files. The code becomes incredibly slow if there is not sufficient memory is available.

I get an error message when I launch a simulation

Check that .NET framework version 4.6 or later is installed on your system.

Where should I start my layer sequence in the input file?

The results of the calculation should not depend on which material layer the sequence starts, i.e. a cyclic permutation in the material layer sequence should not change the simulation results (if not the case, it means that the convergence factors are not chosen to be accurate enough).

How many periods do I have to define in my QCL input file?

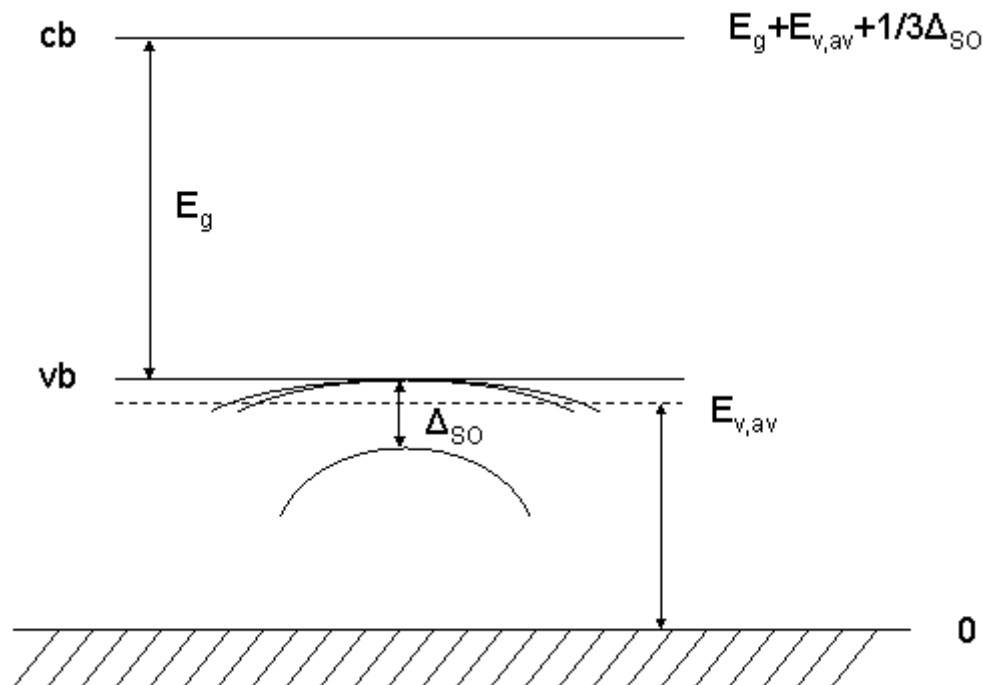
It is sufficient for the standard user to specify only **one** period. A developer (or curious user) might want to have more than one period which can be done manually, i.e. by repeating the well/barrier structure and the doped regions. However, we don't see any physical interest in doing so (except testing the code for consistency). The calculation and convergence will only be longer. In case you want to account for longer coherence length only, the number of periods of coherence can be

increased ($\langle \text{Coherence_length_in_Periods} \rangle$). Concerning the gain, no matter how many periods one simulates, the gain spectra should remain the same. Indeed the material gain is averaged over one period.

How is the conduction band offset defined?

There are two options.

- Option a) Specify conduction band offset (CBO) directly, and then the valence band offset (VBO) is calculated.
- Option b) Specify valence band offset (VBO), and then the conduction band offset (CBO) is calculated.



$E_{v,av}$ is the average energy of heavy hole (hh), light hole (lh) and split-off hole (so).
 Δ_{so} is the spin-orbit split-off energy.

These two different options have different consequences in how the temperature dependence of the bandgap is accounted. Indeed:

- Option a) Specify conduction band offset (CBO) E_c
`<UseConductionBandOffset>yes</UseConductionBandOffset>`

As a consequence, the band offset of the light hole becomes temperature dependent: $E_{hh}(T) = E_c - E_{gap}(T)$

- Option b) Specify valence band offset (VBO) $E_{v,av}$
 The conduction band edge E_c is calculated and depends on temperature.
`<UseConductionBandOffset>no</UseConductionBandOffset>` (default)

$$\begin{aligned} E_{hh} &= E_{v,av} + \frac{1}{3} \Delta_{so} \\ E_c(T) &= E_{v,av} + E_g(T) \end{aligned}$$

$hh\} + E_{\text{gap}}(T) \end{align*}$

Note that the band gap E_{gap} is temperature dependent (Varshni formula), $E_{\text{gap}}(T) = E_{\text{gap}}(T=0 \text{ K}) - \frac{\alpha T^2}{T + \beta}$, where α and β are the Varshni parameters.

How shall I choose the grid spacing "<Spatial_grid_spacing unit=nm>"?

Usually the layers in a QCL have thicknesses of around 1 nm, e. g. 1.3 nm and 1.7 nm. Therefore, in this case, does a grid spacing of 0.2 nm and 0.3 nm make a big difference in the results?

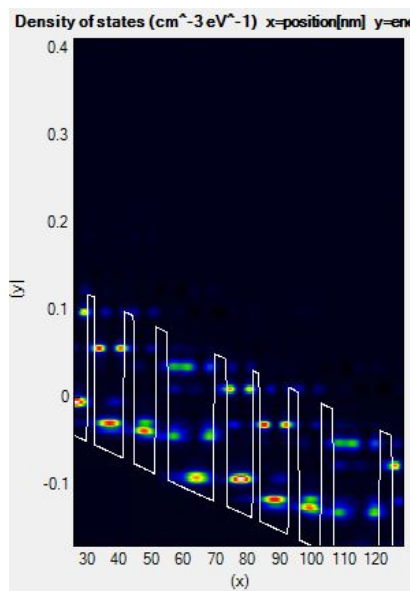
Usually the difference should not be large in this case. However, this should be checked for each structure separately. 0.2 nm to 0.3 nm is a reasonable number for the grid spacing.

What are the meanings of "Check G lesser = 0.922712086509952" and "Check spectral function = 0.917313676654421"?

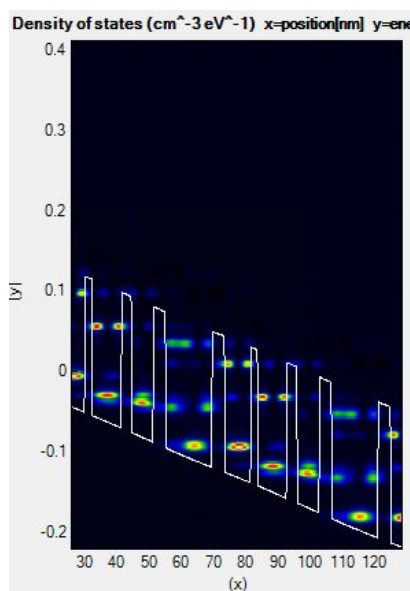
If the final values differ significantly from 1, these are indications that some states are not correctly accounted by the energy grid: either the grid spacing is too coarse, and/or the energy range is not correctly defined (see the role of `Emin_shift` and `E_max_shift` below)

I don't understand the meaning of "Emin_shift" and "E_max_shift" of the energy grid. Can you give an example?

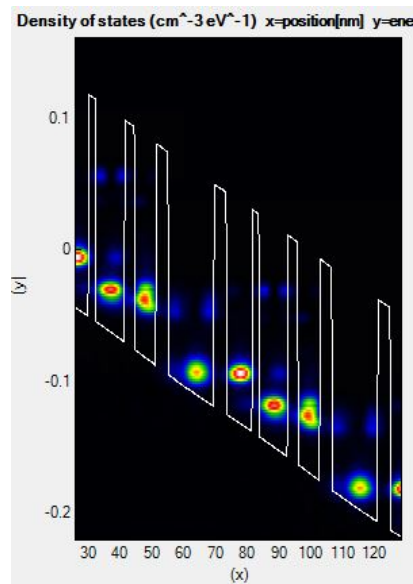
These values are related to the upper and lower limits of the energy range. In the figure below, using `<Emin_shift unit="meV">50</Emin_shift>` and `<Emax_shift unit="meV">0</Emax_shift>` is not a good choice because the lower edge of the energy scale is too high.



The next figure shows the results when we shift `<Emin_shift unit="meV">0</Emin_shift>` and `<Emax_shift unit="meV">0</Emax_shift>`, which gives a better convergence and more accurate results.



Yet, the upper edge goes far beyond the occupied spectral range. Using `<Emin_shift unit="meV">0</Emin_shift>` and `<Emax_shift unit="meV">-250</Emax_shift>` in the figure below leads to faster calculation times while correctly accounting for both the lower and the upper edge of the energy scale.



How is the number of periods considered when calculating the gain?

There is no need to multiply the gain by the period number. The gain is given in (cm^{-1}), not in $1/\text{period}$. The gain is averaged over a single period. But anyway the gain (in cm^{-1}) should not depend on the number of periods considered.

Q2: But for QCLs, there are many periods in the active region. However, in the input file, no period number is considered. In real QCL, the gain should be proportional to the period number. Thus, I don't understand how this number is considered in the NEGF.

A2: The gain per unit length, usually given in cm^{-1} (Wikipedia: [Gain \(laser\)](#)), does not depend on the number of periods. But I guess you refer instead to the round-trip gain in the cavity (Wikipedia: [Round-trip gain](#)), given by $g \cdot 2 \cdot \text{length}(\text{of active region})$.

What value shall I choose for "<Lateral_motion><!--%% Lateral energy spacing %%--%%><Value unit=%%"%%meV%%"%%>"?

It has to be smaller than the linewidth of the states (that you can see on the 2D DOS plot) but the smaller this value the longer the calculation time. We recommend around 4-5 meV for a THz QCL design and around 10 meV for a mid-infrared design. A value up to 40 meV should be sufficient for typical mid-infrared QCLs (and the calculation is much faster than for 10 meV). Large values result in an overestimate of the broadening, which in turn helps the convergence with coarse energy grid. But it is not so accurate. In fact it is needed to reduce simultaneously the energy grid and the lateral energy spacing.

The lateral energy spacing is related to the discretization of the energy dispersion in the directions perpendicular to the growth axis, in which the electrons are free to move (subbands).

My results do not agree with experiment.

A discrepancy in the used **effective masses** could be an explanation. On the other hand, the **interface roughness** parameters are important parameters. The values given in the paper of A. Wacker seem to be taken in order to fit the experimental data. However, their model for elastic scattering is simplified (with no in-plane dependence for the scattering processes), so it might be that the actual parameters for interface roughness are different in reality.

"Current-Density_vs_position.dat": As all the layer are electrically connected in series, would you please advise why the current density varies at different places?

A: The NEGF calculation is done within the mode space basis that depends on the axial cutoff (`<Energy_Range_Axial>`). As a consequence, back in real space, the current is not necessarily conserved, especially through high barriers. By increasing the axial cutoff, the current should be better conserved.

Q: Regarding model #1, i.e. `<Model_Temperature_for_Screening> 1`
`</Model_Temperature_for_Screening>`. If in model #1, the T_{eff} is chosen as a value that matches the "effective electronic temperature" as shown in the output file `Current-miscellaneous.txt`, does it mean that the offset temperature in model #1 is more accurate, since the effective electronic temperature is calculated iteratively?

A: Yes, exactly.

In the gain simulation, some energy points are missing.

In the input file, the energy interval, i.e. the energy spacing between two photon energies (`<dE_Phot unit="meV">` or `<dE_Phot_Self_Consistent unit="meV">`), has been set to be 2 meV. However, for the final results, the interval is 4 meV.

The energy interval for the gain calculation will always be at least the energy grid spacing `<Energy_grid_spacing unit="meV">`.

The simulation results contain a "Gain/" folder but it is empty.

Note that the gain output is only done for the voltages specified in the input file.

```
<!-- Calculate gain only between the following values of
      potential drop per period in order to save CPU time -->
<Vmin unit="mV"> 160 </Vmin>
```

```
<Vmax unit="mV"> 400 </Vmax>
```

What is the difference between the output files "RealSpaceModes.dat" and "Wannier-Stark_levels.dat"? How are they related to the usual approach of calculating the eigenstates in a given conduction band edge profile (single-band Schrödinger equation)?

The output file `Wannier-Stark_levels.dat` gives the usual eigenstates of the conduction band profile for the periodic heterostructure, by solving the single-band Schrödinger equation (with/without nonparabolicity). The output file `RealSpaceModes.dat` gives the position eigenstates within the subspace obtained after applying the axial cut-off energy. These position eigenstates are used as a basis in the NEGF calculation. Note that these states depend on the axial cut-off energy: the larger the axial energy cut-off is, the more localized they are.

The self-consistent gain and semi-classical gain show maximum at different photon energies. Which one to trust more?

The semi-classical calculation is made in the Wannier-Stark states, so it is expected to give maximum photon energy around the same energy as the Wannier-Stark transition energies (though it can be slightly offset as in general multiple peaks are added).

On the other hand, the self-consistent (fully quantum) simulation does not consider any preferred basis and accounts more accurately on broadening. If the broadening (induced by scattering processes are small), semi-classical and self-consistent calculations should give the same result. However, as broadening becomes important, there will be a red shift with respect to the bare transition energies. This shift will depend on the scattering processes. So then the question of which one to trust more is also related to the question whether the parameters for scattering (interface roughness, Coulomb scattering...) matches the reality. And it should be kept in mind there are some underlying assumptions in the NEGF model (in particular the self-consistent Born approximation) which could lead to deviation with respect to reality (such as an overestimate of the red-shifting effect of transition energy with broadening).

At zero bias, when the current asymptotically approaches 0, the current convergence factor does not converge to zero. Is this ok?

When the current approaches 0, it is indeed normal that the current convergence factor does not go to 0. In this case, the convergence should be checked accordingly to the other convergence factor which is based on the lesser Green's function.

From:

<https://nextnano-docu.northeurope.cloudapp.azure.com/dokuwiki/> - **nextnano.NEGF -
Software for
Quantum Transport**

Permanent link:

<https://nextnano-docu.northeurope.cloudapp.azure.com/dokuwiki/doku.php?id=qcl:faq>



Last update: **2022/10/12 10:27**